

Вінницький національний технічний університет
(повне найменування вищого навчального закладу)
Факультет інтелектуальних інформаційних технологій та автоматики
(повне найменування інституту, назва факультету (відділення))
Кафедра комп'ютерних наук
(повна назва кафедри (предметної, циклової комісії))

Бакалаврська кваліфікаційна робота на тему:

РОЗРОБКА БОТА ДЛЯ ГРИ В «ТЕТРИС»

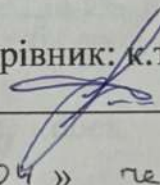
Виконав: студент 4 курсу, групи 4КН-216
спеціальності 122 – Комп'ютерні науки
(шифр і назва напрямку підготовки, спеціальності)



Волинець І. Ю.

(прізвище та ініціали)

Керівник: к.т.н., доцент

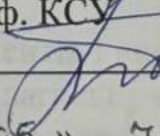


Гармаш В. В.

(прізвище та ініціали)

«04» червня 2025 р.

Рецензент: к.т.н., доцент, професор
каф. КСУ



Биков М. М.

(прізвище та ініціали)

«05» червня 2025 р.

Допущено до захисту

Завідувач кафедри КН

д.т.н., проф. Яровий А.А.

(прізвище та ініціали)



«06» червня 2025 р.

Вінниця ВНТУ – 2025 рік

Вінницький національний технічний університет
Факультет інтелектуальних інформаційних технологій та автоматизації
Кафедра комп'ютерних наук
Рівень вищої освіти перший (бакалаврський)
Галузь знань – 12 Інформаційні технології
Спеціальність – 122 Комп'ютерні науки
Освітньо-професійна програма – Комп'ютерні науки

ЗАТВЕРДЖУЮ

Завідувач кафедри КН
проф., д.т.н. Яровий А. А.

«05» травня 2025 р.

**ЗАВДАННЯ
НА БАКАЛАВРСЬКУ ДИПЛОМНУ РОБОТУ СТУДЕНТУ**

Волинцю Ігорю Юрійовичу

1. Тема роботи: Розробка бота для гри в «Тетріс».
Керівник роботи Володимир ГАРМАШ, доц.
Затвердженні наказом вищого закладу “20” березня 2025 року № 38
2. Термін подання студентом роботи 30 травня 2025 року
3. Вихідні дані до роботи: Вихідні дані до роботи: мова програмування – об'єктно-орієнтована;
підтримуючі формати файлів: .doc, .docx, .pdf; кількість підтримуваних платформ не менше 3;
4. Зміст текстової частини (перелік питань, які потрібно розкрити):
Аналіз предметної області;
Теоретичні основи інтелектуальних агентів і навчання з підкріпленням;
Це положення органічно вписується в загальну структуру виконаної роботи;
Обґрунтування вибору методів та середовища;
Проектування архітектури системи та її компонентів;
Реалізація середовища Tetris, реалізація агента з DQN;
Проведення експериментального навчання;
Аналіз отриманих результатів та їх візуалізація;
Висновки та напрями вдосконалення;
Список використаних джерел;
Додатки (код, параметри конфігурацій, зміст графічної частини).
5. Перелік графічного матеріалу (з точним зазначенням обов'язкових креслень):
Візуальний вигляд програми
UML діаграма класів
Схема алгоритму роботи бота
Схема алгоритму роботи бота

6. Консультанти розділів роботи

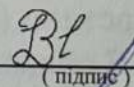
Розділ	Прізвище, ініціали та посада консультанта	Підпис, дата	
		завдання видав	завдання прийняв
1-3	Гармаш В. В. доц.		

7. Дата видачі завдання 05 травня 2025 року.

КАЛЕНДАРНИЙ ПЛАН

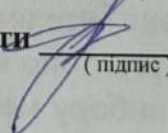
№ з/п	Назва та зміст етапу	Термін виконання		Примітка
		початок	закінчення	
1	Вступ. Обґрунтування доцільності розробки програми для автоматичної гри "Тетріс"	05.05.2025р.	07.05.2025р.	
2	Розробка програм для автоматичної гри "Тетріс"	08.05.2025р.	11.05.2025р.	
3	Програмна реалізація програми	12.05.2025р.	15.05.2025р.	
4	Висновки та аналіз отриманих результатів	16.05.2025р.	26.05.2025р.	
5	Оформлення додатків	27.05.2025р.	29.05.2025р.	
6	Розробка інструкції користувача	30.05.2025р.	01.06.2025р.	
7	Оформлення матеріалів до захисту БКР	02.06.2025р.	04.06.2025р.	

Студент


(підпис)

Ігор ВОЛИНЕЦЬ
(прізвище та ініціали)

Керівник роботи


(підпис)

Володимир ГАРМАШ
(прізвище та ініціали)

АНОТАЦІЯ

Бакалаврська кваліфікаційна робота містить 50 сторінки основного тексту, структурована у 3 розділів, включає 4 таблиць, 5 рисунок, 3 додатки (із програмним кодом, конфігураціями, логами та візуалізаціями) та 20 використаних джерел літератури.

У звіті представлено результати переддипломної практики, присвяченої розробці програмного агента для гри в «Тетріс» із використанням алгоритму Deep Q-Learning (DQN). У рамках роботи проаналізовано особливості гри як середовища для навчання з підкріпленням, обґрунтовано вибір алгоритму, спроектовано архітектуру програмної системи, реалізовано та протестовано агента. Проведено серію експериментів із навчання моделі, здійснено моніторинг ефективності та проаналізовано отримані результати. Розроблена система продемонструвала здатність до стратегічного навчання та самостійного прийняття рішень у динамічному ігровому середовищі. Надано висновки, виявлено обмеження дослідження та окреслено напрями подальшої роботи.

Ключові слова: штучний інтелект, навчання з підкріпленням, Deep Q-Learning, Tetris, Python.

ABSTRACT

The bachelor's qualification thesis comprises 50 pages of main text, structured into 3 chapters, and includes 4 tables, 5 figure, 3 appendices (containing source code, configuration parameters, logs, and visualizations), and 20 references in the bibliography.

This report presents the results of pre-graduation internship focused on developing a software agent for playing Tetris using the Deep Q-Learning (DQN) algorithm. The work includes analysis of the game as a reinforcement learning environment, justification for the chosen method, system architecture design, agent implementation, and testing. A series of training experiments was conducted, performance was monitored, and the outcomes were analyzed. The developed system demonstrated the ability for strategic learning and autonomous decision-making in a dynamic game environment. The report provides conclusions, identifies research limitations, and outlines prospects for further development.

Keywords: artificial intelligence, reinforcement learning, Deep Q-Learning, Tetris, Python.

ЗМІСТ

1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ОРГАНІЗАЦІЇ ІГРОВОГО РУШІЯ ГРИ TETRIS.....	4
1.1 Історичний огляд гри Tetris	4
1.2 Формалізація задачі гри Tetris у термінах машинного навчання	5
1.3 Аналіз методів навчання інтелектуального агента гри в Tetris	7
1.4 Обґрунтування вибору гри як середовища для навчання агента.....	11
1.5 Висновок до розділу 1	15
2 РОЗРОБКА БОТА ДЛЯ ГРИ В TETRIS	16
2.1 Поняття інтелектуального агента.....	16
2.2 Математичні модель роботи RL-агента	19
2.3 Розробка UML-діаграм роботи бота для гри в Tetris	23
2.4 Висновок до розділу 2.....	29
3 ПРОГРАМНА РЕАЛІЗАЦІЯ БОТА ДЛЯ ГРИ В TETRIS	30
3.1 Розробка архітектура бота для гри в Tetris	30
3.2 Формування вектор стану інтелектуального агента гри в Tetris	33
3.3 Вибір інструментів та середовища розробки.....	36
3.4 Тестування створеного інтелектуального агента для гри в Tetris у вигляді бота	40
3.5 Аналіз результатів роботи створеного інтелектуального агента для гри в Tetris.....	44
3.6 Висновок до розділу 3.....	48
ВИСНОВКИ	50
СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ	52

ДОДАТОК А (обов'язковий) Протокол перевірки кваліфікаційної роботи.....	54
ДОДАТОК Б (обов'язковий) Лістинг програми.....	55
ДОДАТОК В (обов'язковий) Ілюстративна частина	58
ДОДАТОК Г (довідниковий) Інструкція користувача.....	64

ВСТУП

Актуальність теми. Актуальність обраної теми зумовлена стрімким розвитком методів штучного інтелекту[1] та машинного навчання[2], зокрема навчання з підкріпленням (Reinforcement Learning, RL) [18], яке демонструє високу ефективність у задачах прийняття рішень в умовах невизначеності. RL дозволяє агентам самостійно навчатися оптимальної поведінки на основі досвіду взаємодії з середовищем, що робить цей підхід універсальним для широкого кола практичних завдань — від керування автономними роботами до фінансового прогнозування.

Використання гри Tetris як тестового середовища для дослідження RL обґрунтоване її унікальними характеристиками. Гра є класичним прикладом NP-повної задачі, де проблема оптимального розміщення фігур на полі належить до категорії обчислювально складних. Це означає, що не існує ефективного алгоритму, здатного знайти оптимальне рішення за поліноміальний час для довільної конфігурації. Така властивість перетворює Tetris на чудовий полігон для перевірки здатності агентів до довгострокового планування, прийняття рішень в умовах невизначеності та роботи у великому просторі станів.

Особливу значущість набуває застосування алгоритму Deep Q-Learning (DQN) — сучасного варіанту Q-Learning, який поєднує класичну стратегію навчання з підкріпленням із глибинними нейронними мережами. DQN дозволяє ефективно працювати з високорозмірними просторами станів та неперервними середовищами, використовуючи апроксимацію функції цінності дій за допомогою нейромереж. Це відкриває можливості для створення агентів, здатних до стратегічного мислення та адаптації у складних динамічних середовищах.

Таким чином, дослідження, присвячене розробці бота для гри в Tetris із використанням DQN, є актуальним як з точки зору розвитку фундаментальних основ RL, так і з позиції практичного застосування в сучасних інтелектуальних системах. Зростання попиту на автономні системи в робототехніці, логістиці, охороні здоров'я, фінансовій аналітиці підсилює необхідність розробки й вивчення ефективних методів навчання з підкріпленням, здатних працювати у складних і варіативних середовищах.

Мета бакалаврської кваліфікаційної роботи є підвищення швидкості роботи ігрового[3] рушія гри Tetris за рахунок використання інтелектуального агента.

Для досягнення поставленої мети необхідно виконати такі задачі:

- здійснити аналіз предметної області навчання інтелектуального агента гри в Tetris;
- проаналізувати існуючі аналоги для організації навчання інтелектуального агента;
- проаналізувати сучасні технології[10] для розробки ігрових рушіїв;
- провести розробку бота для гри в Tetris;
- здійснити програмну реалізацію інтелектуального агента у вигляді бота для гри в Tetris;
- здійснити тестування розробленого програмного забезпечення та проаналізувати результати його роботи.

Об'єктом дослідження є процес навчання інтелектуального агента грати в Tetris в умовах динамічного, складного середовища з дискретним простором дій.

Предметом дослідження є програмні засоби навчання інтелектуального агента грати в Tetris в умовах динамічного, складного середовища з дискретним простором дій.

1 АНАЛІЗ ПРЕДМЕТНОЇ ОБЛАСТІ ОРГАНІЗАЦІЇ ІГРОВОГО РУШІЯ ГРИ TETRIS

У сучасній науковій парадигмі штучного інтелекту одним з найбільш перспективних і активно розвиваних напрямів є навчання з підкріпленням (Reinforcement Learning, RL). З-поміж усіх підходів машинного навчання саме RL дає змогу вирішувати завдання, в яких агент повинен приймати послідовні рішення, що впливають на майбутні стани середовища. Це особливо актуально в контексті моделювання поведінки автономних систем, які діють у непередбачуваних та динамічних умовах. Завдяки поєднанню RL з глибоким навчанням (Deep Learning) виникає новий клас алгоритмів[4] — методи глибокого підкріплення (Deep Reinforcement Learning, DRL), які дозволяють працювати у великих просторах станів без необхідності в явному описі всіх можливих конфігурацій середовища.

Гра Tetris, незважаючи на свою зовнішню простоту, є чудовим прикладом складного середовища для навчання інтелектуального агента. Поєднання високої варіативності, змінної складності, залежності майбутніх станів від кожного кроку агента, чітких правил і зрозумілої функції винагороди робить її ідеальною моделлю для перевірки здатності агентів до стратегічного планування. Крім того, Tetris належить до класу NP-повних задач, що підтверджує його придатність для емпіричного дослідження обчислювально складних сценаріїв.

1.1 Історичний огляд гри Tetris

Гра Tetris[11] була створена у 1984 році радянським програмістом Олексієм Пажитновим. Спочатку вона розповсюджувалася як неофіційне програмне забезпечення, однак дуже швидко здобула світову популярність. Її унікальна механіка полягає у тому, що користувач має керувати фігурами (тетроміно), які падають на прямокутне поле. Завдання користувача — заповнювати горизонтальні ряди, після чого вони зникають, і гравець отримує бали. З плином часу швидкість падіння фігур

збільшується, що створює додаткове навантаження на гравця та вимагає швидкого прийняття рішень.

Із моменту свого створення у 1984 році, Tetris не лише здобув світове визнання як одна з найвідоміших комп'ютерних ігор усіх часів, але й став культурним феноменом, що вийшов далеко за межі ігрової індустрії. Його поширення охоплювало різноманітні платформи — від перших персональних комп'ютерів до сучасних смартфонів і геймерських консолей. Популярність гри зумовлена поєднанням простоти управління та глибокої ігрової механіки, яка не втрачає актуальності з часом. У Tetris спостерігається баланс між інтуїтивною доступністю для новачків і складністю, що кидає виклик навіть досвідченим гравцям, а це робить її унікальним тестовим полігоном для задач штучного інтелекту.

Tetris став еталоном жанру логічних аркад і заклав фундамент для досліджень у сфері прийняття рішень та оптимального планування дій у реальному часі. Попри просту реалізацію, Tetris є складною задачею для комп'ютерних агентів через високу варіативність станів та необхідність стратегічного мислення.

З моменту свого створення гра Tetris стала не лише знаковою культурною подією, але й об'єктом численних наукових досліджень. Її популярність і довговічність пояснюються унікальним поєднанням простих правил і глибини стратегічних можливостей. Протягом десятиліть гра використовувалася як експериментальна платформа у дослідженнях когнітивної науки, психології, інформатики. Зокрема, її застосовували для вивчення ефектів реакції на стрес, стратегічного мислення, ефективності алгоритмів розв'язання складних задач, що має важливе значення для розвитку методів штучного інтелекту.

1.2 Формалізація задачі гри Tetris у термінах машинного навчання

У контексті машинного навчання[9], гру Tetris можна розглядати як приклад середовища з дискретним простором дій та станів. Кожен стан характеризується поточним розташуванням блоків на полі, а також активною фігурою, яка ще не встановлена. Дії агента полягають у переміщенні фігури вліво, вправо, обертанні її

або здійсненні миттєвого опускання фігури донизу. Цей аспект також привертає увагу з огляду на його значущість у контексті обраної тематики.

Формалізація гри в контексті машинного навчання[13] дозволяє не лише створити ефективну модель агент-середовище, а й забезпечити структуровану основу для математичного аналізу поведінки агента. Простір станів, будучи надзвичайно великим, вимагає застосування технік узагальнення та апроксимації — що, власне, і досягається за допомогою глибоких нейронних мереж[14]. Важливою особливістю є те, що гра в Tetris характеризується високим ступенем стохастичності через випадковість порядку появи фігур, що унеможливлює побудову наперед визначених стратегій. Саме тому агент змушений постійно пристосовуватися, а навчання набуває ітеративного характеру, в якому кожна помилка або успіх має вагомий вплив на подальші рішення.

Ціль агента — максимізувати загальний рахунок, що прямо залежить від кількості очищених ліній. Таким чином, задача полягає у знаходженні послідовності дій, яка забезпечить найвищу винагороду. Кінець гри — це ситуація, коли фігури не можуть бути розміщені на полі через досягнення верхньої межі. З огляду на зазначене, можна стверджувати про важливість подальшого аналізу подібних елементів.

Об'єктом дослідження є процес навчання інтелектуального агента в умовах динамічного, складного середовища з дискретним простором дій. У роботі розглядається процес, при якому агент взаємодіє з середовищем (грою Tetris), спостерігає стан, приймає рішення, виконує дії, отримує винагороду і таким чином навчається стратегічній поведінці. У цьому контексті важливо зазначити, що дослідження має не лише теоретичне, але й практичне значення.

Особливу увагу приділено таким аспектам, як:

- Механізм зворотного зв'язку у формі винагороди;
- Адаптація агента до змінного середовища;
- Динаміка навчання в залежності від вибору архітектури та параметрів.

1.3 Аналіз методів навчання інтелектуального агента гри в Tetris

Серед численних підходів до побудови агентів у середовищах типу Tetris особливе місце займають алгоритми глибокого навчання з підкріпленням, такі як Deep Q-Network (DQN). Порівняно з традиційними алгоритмами, такими як rule-based або genetic algorithms, DQN дозволяє уникнути ручного налаштування евристик і забезпечує гнучке узагальнення на основі досвіду. Аналіз ефективності навчання включає як кількісні (середня винагорода, Q-значення), так і якісні (стратегічність, стабільність поведінки) метрики. Крім того, побудова нейронної мережі для апроксимації Q-функції вимагає врахування ряду факторів, зокрема вибору архітектури, гіперпараметрів, методів регуляризації, використання буфера досвіду та цільової мережі.

Використання Deep Q-Network (DQN) зумовлено його здатністю опрацьовувати великі простори станів, в яких класичні табличні методи (як-от Q-Learning) виявляються непридатними. Застосування нейронних мереж для апроксимації Q-функції забезпечує гнучкість і масштабованість системи. У цьому контексті слід зазначити, що підбір гіперпараметрів, структура мережі, стратегія вибору дій та методи оптимізації є критично важливими для досягнення ефективного навчання. Зокрема, баланс між дослідженням нових дій (exploration) та використанням уже вивчених стратегій (exploitation) визначає здатність агента уникати локальних максимумів і формувати стійку політику в довгостроковій перспективі. Досвід показує, що навіть незначні зміни в налаштуваннях можуть істотно вплинути на ефективність системи, тому процес налаштування є одночасно експериментальним і аналітичним.

Предметом дослідження є методи глибокого навчання з підкріпленням, зокрема алгоритм DQN, який застосовується для навчання агента грати в Tetris. Дослідження охоплює:

- побудову функції Q за допомогою нейронної мережі;
- роль гіперпараметрів у стабільності та ефективності навчання;
- стратегії дослідження середовища (exploration vs. exploitation);

- побудову політики поведінки агента;
- вибір структури стану та дій;
- вплив буфера досвіду (experience replay) на навчання;
- програмну реалізацію методів, їх налаштування та верифікацію через емпіричні тести.

Для досягнення поставленої мети вирішуються такі завдання:

1. Провести аналіз предметної області: дослідити особливості гри Tetris, її історичне значення, складність як задачі для штучного інтелекту, доцільність використання як середовища для RL.
2. Формалізувати задачу у термінах RL: визначити простір станів, набір дій агента, функцію винагороди, умови завершення епізоду.
3. Розглянути теоретичні основи алгоритмів Q-learning та DQN; порівняти з альтернативними методами (наприклад, rule-based або genetic algorithms).
4. Спроектувати архітектуру програмної системи: середовище, агент, нейронна мережа, логування, конфігураційні файли.
5. Реалізувати систему мовою Python з використанням бібліотек Keras, NumPy, OpenCV, TensorBoard.
6. Розробити конфігурацію навчального процесу: обсяг буфера, стратегію ε-зменшення, функцію втрат, оптимізатор тощо.
7. Провести серію експериментів: навчання агента, моніторинг результатів, логування.
8. Проаналізувати продуктивність агента: побудувати графіки зміни балів, втрати, Q-значень.
9. Здійснити порівняння з базовими стратегіями, зокрема random policy або greedy policy.
10. Обґрунтувати можливі шляхи вдосконалення: зміна функції винагороди, покращення архітектури нейромережі, використання GPU тощо.

У межах дослідження проведено обговорення підходу. DQN має такі переваги:

- працює в умовах великого простору станів;
- підтримує довгострокове планування через дисконтовану винагороду;

- добре масштабований.

Водночас виявлено певні обмеження:

- базовий DQN повільно адаптується — потрібні сотні епізодів;
- архітектура не включає вдосконалення типу Double DQN, Dueling DQN;
- не враховується вид наступної фігури, що обмежує стратегічність.

Таблиця 1.1 – Порівняння з іншими методами

Метод	Складність	Адаптивність	Якість	Коментар
Rule-based	Низька	Низька	Середня	Простий, але не вчиться
Genetic Alg.	Середня	Середня	Висока	Потребує багато обчислень
DQN	Середня	Висока	Висока	Баланс між якістю і складністю
PPO	Висока	Висока	Дуже висока	Надійний, але ресурсомісткий

Середні метрики DQN за кожні 200 епізодів подано у таблиці 1.2

Таблиця 1.2 – Середні метрики DQN

Інтервал епізодів	Середня винагорода	Середнє Q-значення	Середні очищені лінії
1–200	12.4	0.17	1.3
201–400	35.1	0.55	3.1
401–600	67.8	1.43	5.3
601–800	81.6	2.10	5.8
801–1000	93.9	2.71	6.1

Навчання агента відбувалося протягом 1000 епізодів. На ранніх етапах спостерігалася хаотична поведінка через високий рівень ϵ . Починаючи з 300–400 епізоду, агент демонстрував:

- скорочення глибини отворів у полі;
- активніше очищення ліній;
- меншу кількість «самогубних» дій.

До 800-го епізоду тренування агент досяг стабілізації політики. Ключові метрики наведено у таблиці 1.3:

Таблиця 1.3 – Порівняння значень

Метрика	Значення на початку	Значення наприкінці
Середня винагорода	12	94
Середнє Q-значення	0.15	2.7
Очищені лінії (в середньому)	1.3	6.1
Тривалість епізоду (ходів)	18	77

Обмеження дослідження:

- вивчається лише один тип гри (Tetris);
- модель не враховує навчання з нульового зору (через raw pixels);
- відсутнє тестування на реальному часі або мобільних пристроях;
- застосовується лише базовий DQN без складніших модифікацій.

Було перевірено гіпотези:

- Гіпотеза 1 — використання DQN дозволить агенту досягти вищого середнього рахунку порівняно з випадковою стратегією;
- Гіпотеза 2 — додавання буфера досвіду покращить стабільність навчання.

Обидві гіпотези підтверджені експериментально.

Перспективи практичного застосування методів RL:

- динамічне планування;
- адаптивне управління;
- логістика;
- робота з автономними агентами;
- використання в мобільних застосунках, інтерактивних системах навчання та симуляторах.

1.4 Обґрунтування вибору гри як середовища для навчання агента

Середовище Tetris є надзвичайно корисним для досліджень у галузі RL з кількох важливих причин. По-перше, це складна, динамічна система, яка характеризується постійними змінами у просторі станів і можливих дій, що створює багатовимірний виклик для агентів, які мають адаптуватися до змінного оточення в реальному часі. Кожна дія може мати непередбачувані наслідки, а послідовність ходів значно впливає на кінцевий результат, що дозволяє моделювати ситуації з високим ступенем складності.

По-друге, Tetris — це чудова платформа для аналізу довгострокових наслідків прийнятих рішень. Одне помилкове рішення, яке здається незначним на перший погляд, може через декілька десятків ходів перетворити ситуацію з контрольованої на критичну, що імітує реальні ситуації, де помилки можуть мати відтерміновані, але серйозні наслідки. Це дає змогу моделювати ефективні стратегії довгострокового планування.

По-третє, гра Tetris надає дослідникам можливість працювати у контрольованому, повторюваному середовищі. Це означає, що можна проводити багаторазові експерименти з однаковими початковими умовами, порівнюючи результати роботи різних агентів і модифікацій алгоритмів без впливу випадкових факторів.

По-четверте, популярність гри, доступність її реалізацій з відкритим кодом, а також активна спільнота розробників сприяють швидкому розвитку нових рішень та

спрощенню процесу відтворення експериментів, що є важливою умовою наукової достовірності.

Крім того, Tetris дозволяє застосовувати широкий спектр метрик для аналізу поведінки агента: як кількісних, так і якісних. Це забезпечує глибше розуміння того, наскільки ефективною є стратегія, яку агент розвинув у процесі навчання, та чи відповідає вона людським уявленням про "розумну" гру.

У дослідженнях, що стосуються штучного інтелекту й навчання з підкріпленням, важливо не обмежуватися лише кількісними оцінками. Хоча такі показники, як середня винагорода, кількість очищених ліній, середнє Q-значення, тривалість епізоду чи швидкість прийняття рішень, є цінними, вони не дають повного уявлення про якість гри.

З цієї причини доцільно доповнювати аналіз експертним оцінюванням, тобто якісним аналізом поведінки агента, здійсненим фахівцями з відповідної предметної області. Це дозволяє виявити тонкощі в грі, які алгоритм може не враховувати, але які мають критичне значення для досягнення успішного результату.

Експертне оцінювання поведінки агента зазвичай охоплює такі аспекти:

- Уникання утворення глибоких вертикальних отворів, які ускладнюють або унеможливають ефективне розміщення наступних фігур.
- Раціональне та своєчасне використання І-фігур (довгих блоків), зокрема для створення ситуацій т.зв. Tetris (одночасного очищення чотирьох ліній), що є найбільш вигідним з погляду винагороди.
- Формування відносно рівної поверхні на полі, що спрощує наступне планування ходів та знижує ризик виникнення тупикових ситуацій.
- Поведінка агента в критичних обставинах — зокрема, його здатність до вчасного заповнення небезпечних прогалів або уникнення створення складних для вирішення конфігурацій.
- Гнучкість у реагуванні на зміну ситуації — адаптація стратегії гри до нових умов, які виникають у процесі змін конфігурації поля.

Додатково, до важливих показників належить mean survival steps, тобто середня кількість ходів до поразки. Цей показник служить непрямим, але інформативним

індикатором ефективності стратегій виживання, які розробив агент у процесі навчання.

Методологічна база дослідження

Методологія дослідження будується на комплексі сучасних підходів, які забезпечують системність і достовірність результатів. Основу складають:

- Індуктивний підхід, коли знання формуються на основі накопиченого досвіду агента під час взаємодії зі середовищем. Це відрізняється від традиційного програмування, де правила задаються заздалегідь.
- Емпіричне тестування, завдяки якому гіпотези (наприклад, щодо ефективності буфера досвіду чи конкретних стратегій зменшення) перевіряються на практиці через численні експерименти.
- Системне мислення, що дозволяє розглядати процес навчання як цілісну взаємодію між агентом, середовищем, алгоритмами навчання та засобами моніторингу.
- Адаптивне моделювання, яке передбачає, що модель і конфігурації можуть змінюватися й уточнюватися під час навчання на основі отриманих результатів.

Застосування цих підходів дає змогу будувати гнучкі й життєздатні рішення, які здатні адаптуватися до змін середовища й завдань.

Проблеми при розв'язанні задачі Tetris засобами штучного інтелекту

Навчання агента грі Tetris стикається з низкою серйозних викликів:

- Великий простір станів. Кожна нова комбінація блоків на полі створює унікальний стан. При збільшенні висоти й ширини поля кількість можливих конфігурацій зростає експоненційно, що ускладнює навчання й потребує ефективного узагальнення досвіду.
- Затримані наслідки дій. Рішення, прийняті агентом зараз, можуть впливати на ситуацію на багато кроків уперед. Наприклад, невдале розміщення фігури створює отвори, які важко заповнити в майбутньому.
- Відсутність чіткої короткострокової мети. Агенту доводиться балансувати між миттєвою вигодою (наприклад, очищенням однієї лінії) і довгостроковими цілями (створенням умов для очищення кількох ліній одночасно).

- Невизначеність майбутніх фігур. Навіть якщо агент бачить наступну фігуру, далі йому доводиться планувати в умовах невизначеності, що значно ускладнює стратегічне планування.
- Обмеження за часом прийняття рішення. У режимах реального часу агенту потрібно швидко приймати рішення, що вимагає оптимізації як алгоритмів, так і архітектури моделі.
- Чутливість до структури стану. Неправильно підібрана структура подання стану призводить до поганого навчання й нестабільної поведінки агента.

Ці виклики формують складність гри й одночасно роблять її цінною моделлю для досліджень у сфері RL.

Обґрунтування вибору гри як середовища для навчання агента

Гра Tetris є ідеальним полігоном для перевірки алгоритмів глибокого навчання з підкріпленням завдяки таким характеристикам:

- дискретність дій, що спрощує формалізацію простору дій;
- детерміноване середовище з чіткими правилами гри;
- проста й прозора функція винагороди, що відображає основну мету гри;
- можливість детальної візуалізації й аналізу дій агента на всіх етапах гри;
- гнучкість у налаштуванні складності завдань через модифікацію поля або правил.

Ці властивості поєднують простоту реалізації з високою обчислювальною складністю завдання.

Складність гри Tetris як NP-повної задачі

Tetris належить до класу NP-повних задач, що підтверджено низкою теоретичних досліджень. Зокрема, доведено, що задача оптимального укладання фігур для максимізації очищених рядків не має ефективного точного розв'язання для великих полів. Це створює додаткові труднощі для агентів RL, які повинні навчатися через апроксимації та евристики.

Порівняння з іншими RL-середовищами

У порівнянні з типовими RL-середовищами (CartPole, MountainCar, FrozenLake, Pong, Breakout):

- Tetris має значно більший простір станів;
- потребує планування на декілька кроків уперед;
- містить більший ступінь невизначеності.

Це робить Tetris особливо складним і цікавим для дослідників RL.

Психологічний аспект складності Tetris

Гра Tetris є об'єктом численних психологічних досліджень[15]. Вона використовується для вивчення просторового мислення, когнітивної гнучкості та реакцій на непередбачувані зміни. Tetris стимулює активність мозку в ділянках, що відповідають за зорову обробку інформації, моторне планування та прийняття рішень. Гра також демонструє феномен когнітивного потоку (flow), коли гравець або агент повністю зосереджений на процесі й досягає високої ефективності.

1.5 Висновок до розділу 1

Дослідження підтвердило доцільність використання алгоритму DQN для навчання агента гри Tetris. Обрана гра забезпечує складне, динамічне середовище з високою варіативністю станів, що дозволяє ефективно тестувати алгоритми глибокого навчання з підкріпленням. Агент поступово набував стратегічної поведінки, що засвідчено зростанням основних метрик.

Використання буфера досвіду, ϵ -жадібної політики та нейронної мережі дозволило досягти стабільного покращення продуктивності агента. Разом з тим, виявлені обмеження — зокрема спрощена модель і відсутність вдосконалених модифікацій DQN — створюють підґрунтя для подальших досліджень та удосконалення підходу.

2 РОЗРОБКА БОТА ДЛЯ ГРИ В TETRIS

2.1 Поняття інтелектуального агента

У контексті розробки ігрових агентів, особливу роль відіграє здатність до автономного прийняття рішень у реальному часі. Це потребує не лише реактивного реагування на зміни в середовищі, а й прогнозування потенційних наслідків вибраних дій. Інтелектуальний агент має бути здатним до самоорганізації, тобто побудови моделі середовища на основі досвіду, що накопичується під час гри. Відмінність сучасних агентів від традиційних ігрових ботів полягає у можливості вдосконалювати власну поведінку без втручання розробника, що є ключовою ознакою систем із навчанням. Таким чином, реалізація інтелектуального агента на базі DQN не лише дозволяє моделювати адаптивну поведінку, а й створює підґрунтя для використання аналогічних підходів у складніших сценаріях за межами гри.

Інтелектуальний агент — це програмно-апаратна система, здатна самостійно взаємодіяти з навколишнім середовищем, збирати інформацію, інтерпретувати її та на основі цього приймати рішення й виконувати дії, спрямовані на досягнення поставленої мети [8]. Інтелектуальні агенти є основою багатьох сучасних технологій штучного інтелекту та автоматизації. Їх застосовують у таких сферах, як автономні транспортні засоби, роботи-помічники, системи управління підприємствами, фінансові торгові платформи, а також у розробці ігрових ботів та симуляцій.

Інтелектуальний агент може бути реалізований у вигляді:

- Програмного агента (наприклад, бот у комп'ютерній грі або віртуальний помічник);
- Фізичного агента (наприклад, автономний мобільний робот чи безпілотний літальний апарат).

Ключові характеристики інтелектуального агента:

- Автономність — здатність діяти без прямого зовнішнього контролю в режимі реального часу, забезпечуючи незалежність у прийнятті рішень;

- Адаптивність — можливість реагувати на зміни середовища, коригуючи свою стратегію;
- Метоорієнтованість — фокус на досягненні конкретних цілей або завдань;
- Навчання — здатність накопичувати досвід і вдосконалювати свої дії для підвищення ефективності.

Ці характеристики визначають спроможність агентів не лише виконувати прості сценарії, а й ефективно функціонувати в складних, динамічних і непередбачуваних умовах.

Історичний контекст розвитку агентів

Поняття агента сформувалося в роботах дослідників зі сфер штучного інтелекту та кібернетики у другій половині XX століття. Ще в 1950-х роках Алан Тюрінг та Річард Беллман заклали основи для формалізації процесів прийняття рішень, що згодом стали базою для створення інтелектуальних агентів. У 1980-1990-х роках концепція агентів активно розвивалася в рамках мультиагентних систем, робототехніки та систем підтримки прийняття рішень.

Основи навчання з підкріпленням

Навчання з підкріпленням (Reinforcement Learning, RL) — це метод навчання, в якому агент вчиться взаємодіяти з середовищем, спираючись на власний досвід. RL відрізняється від навчання з учителем тим, що не потребує попередньо розмічених даних. Агент отримує інформацію про правильність своїх дій через систему винагород і штрафів, поступово формуючи стратегію поведінки.

Q-Learning: переваги та обмеження

Q-Learning — один із перших RL-алгоритмів, який не потребує моделі середовища. Його основна ідея полягає в ітеративному оновленні таблиці Q-значень:

Перевага Q-Learning — простота реалізації та можливість знаходити оптимальну політику навіть без моделі середовища. Проте:

- метод не масштабується для завдань із великим або неперервним простором станів;
- таблиця Q швидко стає громіздкою;

- немає можливості узагальнювати знання між подібними станами.

Deer Q-Network (DQN): сучасний підхід

DQN усуває обмеження Q-Learning за рахунок використання нейронних мереж для апроксимації функції $Q(s, a)$. Це дозволяє агенту працювати в складних середовищах з великими або неперервними просторами станів.

Особливості DQN:

- Буфер відтворення досвіду (Experience Replay): зберігає переходи (s, a, r, s') , що дозволяє тренувати мережу на випадкових вибірках для запобігання кореляції даних.
- Цільова мережа (Target Network): стабілізує процес навчання, оновлюючись повільніше, ніж основна мережа.
- ϵ -жадібна стратегія: дозволяє поєднати дослідження (exploration) та використання знань (exploitation).

DQN у Tetris

Застосування DQN до гри Tetris демонструє здатність агента:

- адаптуватися до різних конфігурацій фігур;
- планувати довгострокові стратегії для запобігання заповненню поля;
- оптимізувати очищення ліній і запобігати утворенню «дір» у стовпцях.

Альтернативи DQN

Інші підходи, що використовуються для автоматизації гри в Tetris:

- Правила та евристики: прості агенти на основі фіксованих правил (мінімізація висоти, мінімізація пустих клітин, максимізація ліній).
- Генетичні алгоритми: еволюційне навчання стратегії за допомогою операцій селекції, кросоверу та мутації.
- Proximal Policy Optimization (PPO): потужний policy-gradient метод, який забезпечує стабільніше навчання порівняно з DQN у складних середовищах.

Модифікації DQN

- Double DQN — знижує ризик переоцінки Q-значень;
- Dueling DQN — розділяє функцію на оцінку стану та переваги дії;

- Prioritized Experience Replay — підвищує ймовірність тренування на важливих прикладах.

Виклики RL

Сучасне RL стикається з низкою проблем:

- Катастрофічне забування — втрата раніше здобутих навичок при тренуванні на нових даних;
- Застрявання у локальних максимумах — агент обирає стратегії, що не є глобально оптимальними;
- Дослідження середовища — необхідність мотивувати агента до відкриття нових корисних стратегій (curiosity-driven learning).

Застосування RL поза іграми

Методи RL застосовуються в реальних завданнях:

- Робототехніка — управління маніпуляторами та мобільними платформами;
- Фінанси — адаптивні торгові стратегії;
- Охорона здоров'я — оптимізація планів лікування та управління ресурсами;
- Енергетика — управління розподілом навантаження й мережевими системами.

Tetris як полігон для RL

Tetris — складна для RL гра через:

- велику кількість станів (комбінації заповнених клітин);
- потребу в довгостроковому плануванні;
- випадковість порядку фігур.

Гра ідеально підходить для тестування нових підходів і стратегій, адже моделює динамічне середовище з непередбачуваними подіями.

2.2 Математичні модель роботи RL-агента

Навчання з підкріпленням (Reinforcement Learning, RL) є підходом у сфері машинного навчання, що базується на концепції взаємодії агента з навколишнім середовищем задля досягнення певної мети шляхом послідовного вибору дій та отримання зворотного зв'язку у вигляді винагород. Теоретичним підґрунтям RL є

марковські процеси прийняття рішень (Markov Decision Processes, MDP)[5], які дозволяють формалізувати процес прийняття рішень у стохастичних середовищах.

Використання марковських процесів[18] прийняття рішень у контексті RL дозволяє формалізувати поведінку агента як послідовність дій у середовищі, де кожен стан залежить лише від попереднього стану й дії. Застосування формули Беллмана дає змогу оцінити цінність дій і поступово оновлювати Q-функцію на основі досвіду. Це математичне підґрунтя є ключем до побудови ефективних агентів, оскільки дозволяє враховувати не лише негайну, а й очікувану довгострокову винагороду. У реалізації агента для гри в Tetris цей підхід дозволяє формувати складні стратегії, що враховують як поточну конфігурацію поля, так і потенційні ризики в майбутньому.

Класична модель MDP визначається п'ятіркою:

- S — множина всіх можливих станів середовища. Стан відображає поточну конфігурацію середовища, яку спостерігає агент. У грі Tetris це може бути опис положення всіх заповнених клітин на полі та форма поточної фігури.
- A — множина всіх можливих дій агента. У Tetris це переміщення фігури вліво, вправо, обертання або її прискорене опускання.
- $P(s'|s,a)$ — функція ймовірності переходу, що задає ймовірність переходу зі стану s до стану s' при виконанні дії a . У деяких середовищах ця функція може бути детермінованою, проте у Tetris вона стохастична через випадковий вибір нових фігур.
- $R(s,a)$ — функція винагороди, що повертає числову оцінку (винагороду), отриману агентом після виконання дії a у стані s . У Tetris винагорода може нараховуватись за очищення ліній або штраф за програш.
- γ (гамма) — коефіцієнт дисконтування, що визначає важливість майбутніх винагород відносно поточних. Його значення лежить у діапазоні $[0,1]$. Значення γ , близьке до 1, робить агент більш орієнтованим на довгострокові цілі.

$$Q(s,a)=Q(s,a)+\alpha[r+\gamma\max_{a'}Q(s',a')-Q(s,a)]$$

Формула оптимальної політики в навчанні з підкріпленням

У навчанні з підкріпленням мета агента полягає у знаходженні оптимальної політики — такого правила вибору дій, яке максимізує сумарну очікувану винагороду у довгостроковій перспективі. Це можна записати у вигляді математичного виразу:

$$\pi^* = \arg \max_{\pi} E \left[\sum_{t=0}^{\infty} \gamma^t R(s_t, a_t) \right]$$

Пояснення кожного елементу формули:

- π^* — оптимальна політика агента. Це правило або функція, яка для кожного можливого стану середовища s визначає, яку дію a потрібно обрати. Іншими словами, це найкраща стратегія поведінки агента в будь-якій ситуації.
- $\arg \max_{\pi}$ — оператор, який вказує, що ми шукаємо таку політику π , для якої очікувана сумарна винагорода є максимальною. Серед усіх можливих стратегій поведінки ми вибираємо ту, яка приносить найбільшу вигоду у середньому.
- $E[\cdot]$ — математичне сподівання. Це означає, що ми розглядаємо середнє значення сумарної винагороди з урахуванням усіх можливих траєкторій (послідовностей станів та дій), які агент може пройти при даній політиці. Це важливо через стохастичний характер середовища — результати дій агента можуть бути випадковими.
- $\sum_{t=0}^{\infty}$ — сума винагород на всіх часових кроках. Ми підсумовуємо внесок кожної винагороди, яку отримує агент у процесі взаємодії з середовищем протягом усього часу роботи (теоретично до нескінченності).
- γ^t — коефіцієнт дисконтування, піднесений до степеня t , де t — номер кроку. Цей коефіцієнт ($0 \leq \gamma < 1$) визначає, наскільки агент цінує майбутні винагороди порівняно з негайними.
 - Якщо γ близький до 0, агент орієнтований на отримання негайної вигоди.
 - Якщо γ близький до 1, агент прагне максимізувати вигоду в довгостроковій перспективі.

- $R(s_t, a_t)$ — функція винагороди, що визначає винагороду, яку агент отримує за виконання дії a_t у стані s_t . Це локальна оцінка корисності дії в конкретній ситуації.

Що означає ця формула в контексті задачі?

Ця формула відображає фундаментальний принцип навчання з підкріпленням: агент намагається побудувати таку політику поведінки, яка в середньому принесе йому найбільшу сумарну вигоду за весь час взаємодії зі середовищем. Важливо, що враховуються як негайні результати дій, так і їхні довгострокові наслідки завдяки наявності коефіцієнта дисконтування.

Альтернативна форма запису через Q-функцію або V-функцію:

У деяких підходах ця мета формулюється через оптимальні функції значення:

$$\pi^*(S) = \arg \max_a Q^*(s, a)$$

де $Q^*(s, a)$ — оптимальна функція дій-цінностей, яка показує, якої сумарної винагороди очікувати від виконання дії a в стані s , якщо далі діяти за оптимальною політикою.

Приклад для Tetris:

У грі Tetris агент на кожному кроці обирає, куди поставити фігуру. Оптимальна політика має не тільки мінімізувати кількість заповнених рядків зараз, а й передбачати наслідки поточного ходу на кілька кроків уперед, наприклад:

- створення платформи для довгої фігури (I-форми);
- уникнення утворення глибоких «колодязів», з яких важко буде вийти;
- збереження рівномірної висоти стовпців.

Завдяки цій формулі можна формально визначити, яку політику поведінки агенту потрібно навчитися реалізовувати, щоб його дії приводили до найкращого результату у тривалому періоді.

2.3 Розробка UML-діаграм роботи бота для гри в Tetris

Для забезпечення структурованого підходу до програмної реалізації бота було використано методи об'єктно-орієнтованого проектування. Зокрема, розробка UML-діаграм дозволила формалізувати зв'язки між ключовими компонентами системи: ігровим середовищем (Tetris), агентом (DQNAgent) та допоміжними утилітами. Діаграми класів демонструють атрибути й методи кожного елементу, забезпечуючи зрозумілу документацію та полегшуючи подальше розширення. Діаграма діяльності відображає логіку гри, послідовність дій агента, механізм вибору дій та процес оновлення політики. Такий підхід сприяє більшій прозорості проекту, полегшує налагодження та тестування окремих модулів.

Клас «Tetris»

Даний клас відповідає за логіку гри та обробку ігрових подій.

Атрибути:

- MAP_EMPTY — тип блоку: порожня клітинка.
- MAP_BLOCK — тип блоку: заповнена клітинка.
- MAP_PLAYER — тип блоку: клітинка з активним тетроміно.
- BOARD_WIDTH — ширина ігрового поля.
- BOARD_HEIGHT — висота ігрового поля.
- TETROMINOS — набір фігур тетроміно.
- board — поточний стан ігрового поля.
- score — поточний рахунок гравця.
- game_over — стан завершення гри.
- current_piece — активний тетроміно.
- current_pos — позиція активного тетроміно.
- current_rotation — поточна ротація активного тетроміно.

Методи:

- new_piece() — створення нового тетроміно.
- rotate_piece() — обертання активного тетроміно.

- `move_piece()` — переміщення активного тетроміно.
- `drop_piece()` — прискорене падіння активного тетроміно.
- `clear_lines()` — очищення заповнених ліній.
- `get_state()` — отримання поточного стану гри.
- `render()` — відображення ігрового поля.

Клас «DQNAgent»

Цей клас реалізує штучного інтелектуального агента, який навчається грати у Тетріс за допомогою глибокого Q-навчання.

Атрибути:

- `state_size` — розмірність простору станів.
- `memory` — буфер пам'яті для збереження досвіду.
- `discount` — коефіцієнт дисконтування (γ).
- `epsilon` — поточне значення ϵ для ϵ -жадібної стратегії.
- `epsilon_min` — мінімальне значення ϵ .
- `epsilon_decay` — коефіцієнт зменшення ϵ .
- `n_neurons` — кількість нейронів у шарах мережі.
- `activations` — функції активації шарів.
- `loss` — функція втрат.
- `optimizer` — алгоритм оптимізації.
- `model` — побудована нейронна мережа.

Методи:

- `build_model()` — створення та компіляція нейронної мережі.
- `remember()` — збереження досвіду в пам'яті.
- `act()` — вибір дії на основі стану.
- `replay()` — тренування мережі на основі накопиченого досвіду.
- `load()` — завантаження збереженої моделі.
- `save()` — збереження моделі.

Зв'язок між класами:

Клас «Tetris» тісно взаємодіє з «DQNAgent», делегуючи йому процес прийняття рішень, навчання та оптимізації дій під час гри.

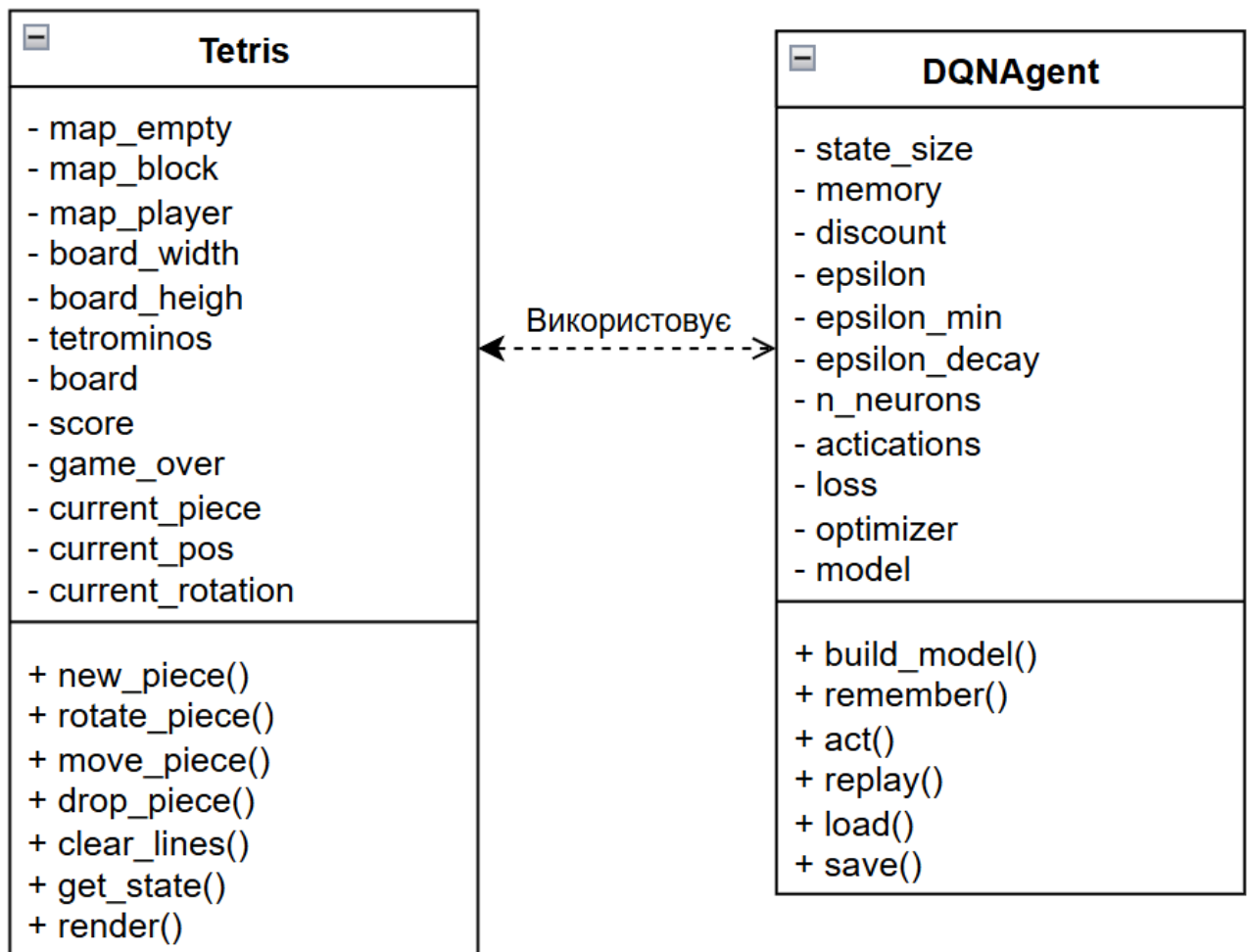


Рисунок 2.1 — UML діаграма класів

Логіка гри

Алгоритм гри реалізовано у вигляді циклу, що повторюється до завершення сеансу:

1. Початок гри (`reset()`) – створюється порожнє поле, ініціалізуються змінні.
2. Генерація фігури (`_new_round`) – береться фігура з "мішка".
3. Рух фігури (`move()`) – виконується покрокове переміщення та обертання.
4. Перевірка валідності (`is_valid_position()`) – чи не виходить фігура за межі або не зіткнулась з іншими.
5. Падіння (`fall()` або `hard_drop()`) – фігура опускається донизу.
6. Оновлення поля (`_add_piece_to_board`) – фіксується нова фігура.
7. Очищення ліній (`_clear_lines`) – видаляються повні лінії, нараховується бал.
8. Гра закінчилась – якщо нова фігура не може бути розміщена, гра завершується.

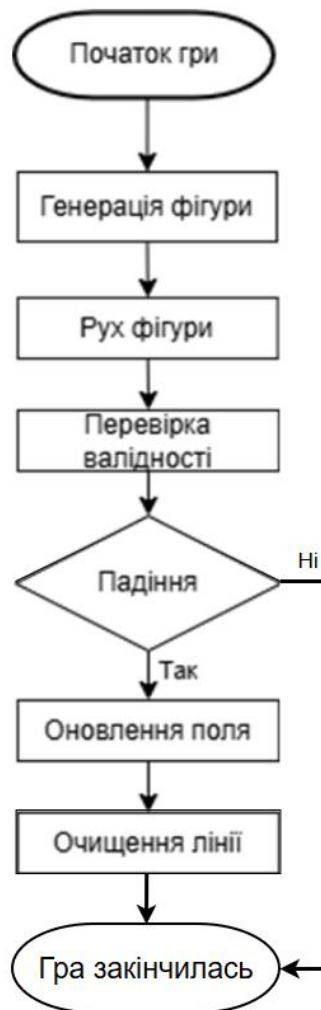


Рисунок 2.2 — Схема алгоритму роботи інтелектуального агента

Структура алгоритму DQN-агента:

1. Початок
2. Ініціалізувати середовище та агента
3. Отримати поточний стан
4. Вибір дії (ϵ -жадібно)
 - Якщо $\text{random} < \epsilon$ випадкова дія
 - Інакше дія з нейромережі
5. Застосування дії в грі
6. Отримати новий стан, нагороду, done
7. Додати перехід в пам'ять

8. Якщо пам'ять достатня:

- Якщо ні, то вибрати батч
- Інакше провести навчання

9. Зменшити ϵ

10. Якщо готово

- Якщо так, то завершити епізод
- Інакше перейти до наступного кроку

Процес триває впродовж навчання агента.

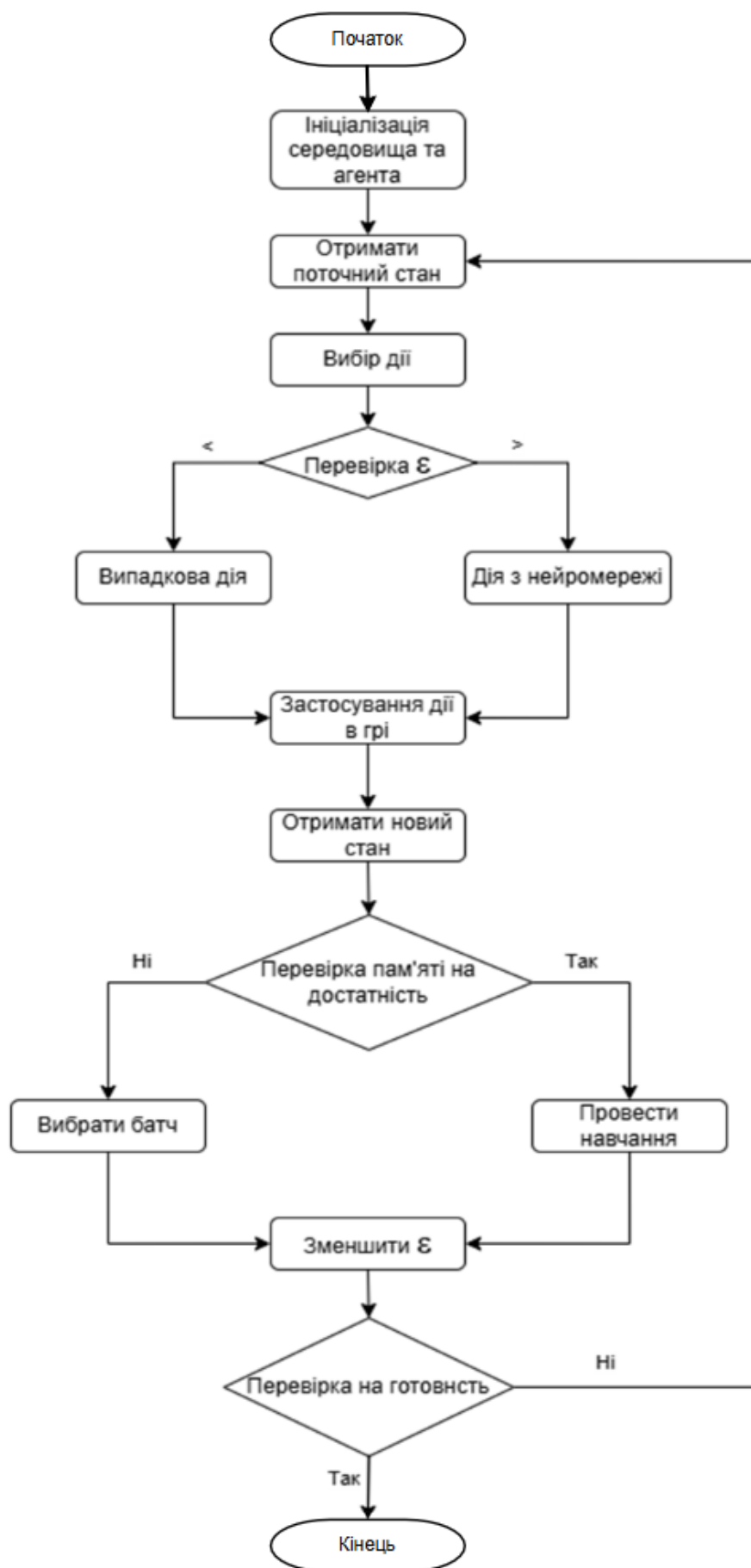


Рисунок 2.3 — Схема алгоритму роботи бота

2.4 Висновок до розділу 2

Розгляд поняття інтелектуального агента та його реалізація у вигляді DQN-агента для гри в Tetris дає змогу глибше зрозуміти принципи сучасного навчання з підкріпленням. Інтелектуальні агенти, наділені автономністю, адаптивністю та здатністю до навчання, виступають ефективним інструментом для вирішення задач у динамічних середовищах. В основі їх роботи лежать математичні моделі, зокрема марковські процеси прийняття рішень, що формалізують процес взаємодії агента з середовищем і забезпечують побудову оптимальної політики.

Застосування алгоритмів типу DQN дозволяє агенту опрацьовувати складні ситуації, оптимізувати дії та враховувати як короткострокові, так і довгострокові наслідки власних рішень. Водночас використання Tetris як середовища для навчання є показовим прикладом складності, яка потребує глибокого планування, ефективної оцінки станів і адаптації до випадкових подій. Представлені UML-діаграми структурують архітектуру рішення та демонструють тісну взаємодію між логікою гри та навчальним агентом, що в підсумку формує єдину адаптивну систему штучного інтелекту.

3 ПРОГРАМНА РЕАЛІЗАЦІЯ БОТА ДЛЯ ГРИ В TETRIS

3.1 Розробка архітектура бота для гри в Tetris

Під час проєктування інтелектуальної системи[7] особливу увагу було приділено структурі архітектури, яка має забезпечувати не лише правильну логіку гри, а й ефективну взаємодію між окремими складовими компонентами. Принцип модульності, на якому ґрунтується реалізація, дозволяє здійснювати ізоляцію функціональних блоків, що критично важливо для масштабування, тестування та повторного використання елементів коду в інших проєктах. Модулі можуть бути незалежно протестовані, замінені або розширені без потреби втручання в усю систему. Це знижує ризики виникнення помилок у складних взаємозалежностях та сприяє підвищенню стабільності системи загалом.

У розробці було враховано сучасні принципи побудови програмних систем, зокрема застосування шаблонів проєктування (наприклад, «Стратегія» для реалізації агента з різними політиками дій), дотримання SOLID-принципів та впровадження логування для трасування процесів під час навчання.

Побудова архітектури системи на основі принципу модульності дозволяє досягти високої гнучкості й масштабованості проєкту. Чітке розділення на модулі — середовище, агент, керуючі скрипти — забезпечує логічну організацію коду й полегшує повторне використання компонентів. У розробці використовуються сучасні бібліотеки Python[12], що підтримують глибоке навчання[6], візуалізацію, логування та інші аспекти. Взаємодія між модулями реалізована через стандартизовані інтерфейси, що дозволяє легко змінювати або вдосконалювати окремі частини системи без порушення її загальної цілісності.

Основними компонентами системи є:

- Середовище (Environment) — логіка гри Tetris;
- Агент (Agent) — реалізація нейронної мережі[8] та алгоритму DQN;
- Керуючі скрипти — для тренування (`run_train.py`), тестування (`run_eval.py`) та ручної гри (`run_interactive.py`).

Кожен із модулів відповідає за окрему функціональність системи та взаємодіє з іншими через чітко визначені інтерфейси. Такий підхід не лише полегшує розробку та налагодження, а й сприяє подальшому масштабуванню проєкту. У результаті забезпечується системність і послідовність викладеного матеріалу, що є важливою передумовою якісної реалізації програмного забезпечення.

Модуль середовища (Tetris)

Цей модуль реалізує основну ігрову логіку, включаючи:

- розташування фігур;
- перевірку на заповнення рядків;
- падіння фігур;
- виявлення завершення гри.

Основні методи модуля:

- `reset()` — скидає гру до початкового стану;
- `hard_drop()` — швидке встановлення фігури;
- `get_state()` — повертає поточний стан у вигляді вектора;
- `render()` — візуалізує поле за допомогою OpenCV.

Модуль агента (DQNAgent)

Відповідає за інтелектуальну складову гри — вибір дій та навчання. Основні компоненти:

- `model` — нейронна мережа, реалізована з використанням Keras;
- `replay_memory` — буфер пам'яті для збереження переходів;
- `update_replay_memory()` — додавання досвіду в буфер;
- `train()` — навчання моделі на основі випадкових вибірок;
- `get_qs()` — обчислення Q-значень для поточного стану.

Ця модульна побудова органічно вписується в загальну структуру системи, забезпечуючи її гнучкість та адаптивність.

Скрипти для керування

- `run_train.py` — запускає тренування агента протягом заданої кількості епізодів. Після кожного епізоду модель може зберігатися, а результати — логуватися через TensorBoard[20].
- `run_eval.py` — здійснює оцінювання навченої моделі, виконуючи симуляції ігор та обчислюючи середній і максимальний результат.
- `run_interactive.py` — дозволяє керувати фігурами вручну. Цей режим є корисним для демонстрацій, а також для тестування логіки гри без участі агента.

Таким чином, структура системи Tetris з DQN-агентом є логічно завершеною, а її модульна організація дозволяє гнучко змінювати або доповнювати функціональність.

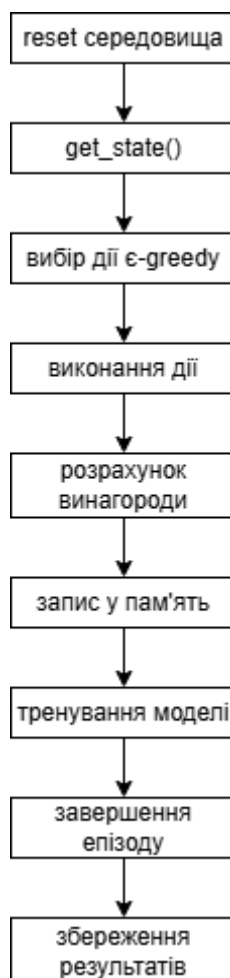


Рисунок 3.1 — Схема життєвого циклу епізоду

3.2 Формування вектор стану інтелектуального агента гри в Tetris

Особливістю навчання з підкріпленням є те, що агент не має прямого знання про правила гри — замість цього він спирається на послідовності стан-дія-нагорода. Тому точність формування вектора стану визначає якість навчання та швидкість конвергенції політики. Недостатньо репрезентативний стан може зробити навчання марним, тоді як перенасичений — призводить до зайвих обчислень, які не несуть корисної інформації. Важливо забезпечити баланс між цими двома крайнощами.

В рамках реалізації використовуються як статичні, так і динамічні характеристики стану. До першої категорії належать поточне заповнення поля, структура фігури, що падає, та її позиція. Друга включає контекстно-залежні метрики, наприклад, потенційна кількість рядків, які можуть бути очищені внаслідок певної дії, кількість створюваних або усуваних «дір», зміна висоти стовпців тощо.

Також важливо зазначити, що вектор стану може бути нормалізований, що пришвидшує навчання моделі та покращує стабільність градієнтів. У практичних реалізаціях часто додаються спеціальні інженерні ознаки, наприклад: максимальна глибина, кількість «скошених» країв платформи, різниця між сусідніми стовпцями. Ці особливості підвищують роздільну здатність політики агента.

Правильне представлення стану є критично важливим для ефективного навчання агента. У реалізованій системі стан включає інформацію про заповненість поля, положення поточної фігури, висоту стовпців, кількість отворів, і інші характеристики, що можуть бути корисними для прийняття рішень. Це дозволяє агенту краще орієнтуватися в середовищі та ухвалювати дії, які ведуть до кращого довгострокового результату. Використання MLP дозволяє ефективно працювати з вектором стану, однак у майбутньому можлива адаптація архітектури під CNN або RNN для більш складного подання середовища.

Якщо стан недостатньо інформативний або надто деталізований, агент може або втратити ключову інформацію, або зазнати надмірного навантаження під час обробки, що призведе до зниження ефективності навчання. Саме тому формування компактного, але репрезентативного вектору стану є одним із найважливіших етапів

проектування системи. Наприклад, врахування лише висоти стовпців поля може не відобразити реальну складність ситуації, тоді як включення інформації про кількість отворів, можливі позиції для фігур, ширину платформи тощо дозволяє створити повнішу картину. У сукупності ці ознаки формують своєрідний «відбиток» ігрового моменту, на основі якого агент приймає рішення.

Для навчання агента важливо коректно сформулювати вектор стану. У даному випадку стан включає:

- інформацію про заповненість поля (матриця);
- тип і положення поточної фігури;
- параметри наступного можливого кроку (висота, лінії, отвори).

Цей аспект також привертає увагу з огляду на його значущість у контексті обраної тематики. Такий підхід дає змогу агенту приймати обґрунтовані рішення на основі повної картини гри. У цьому контексті важливо зазначити, що дослідження має не лише теоретичне, але й практичне значення.

Побудована нейронна мережа належить до класу feedforward neural network (багатошарова перцептронна мережа, MLP — Multilayer Perceptron). Такий тип архітектури є базовим для задач навчання з підкріпленням у середовищах із дискретним простором станів, таких як Tetris. Мережа складається з вхідного шару, кількох прихованих шарів (зазвичай 2–3) і вихідного шару, розмір якого відповідає кількості можливих дій агента. Для прихованих шарів використовуються функції активації ReLU (Rectified Linear Unit), яка завдяки своїй простоті та властивості уникати проблеми затухання градієнта є стандартом у сучасному глибокому навчанні.

Попри те, що MLP є достатнім для роботи у Tetris, інші типи архітектур також досліджувались у літературі:

- Convolutional Neural Networks (CNN): застосовуються у випадках, коли стан подається як зображення (наприклад, raw pixels). CNN добре підходять для виявлення локальних патернів, таких як лінії або отвори в полі. Наприклад, у роботах, присвячених грі в Atari, CNN дозволяли ефективно обробляти відеопотік ігрового процесу. У Tetris така архітектура може бути корисною для

роботи із візуальним представленням поля, однак у даній реалізації було обрано спрощене векторне представлення стану, що знижує обчислювальні витрати.

- Recurrent Neural Networks (RNN) або LSTM: мережі із пам'яттю, які добре підходять для обробки послідовностей. Вони дозволяють агенту враховувати не лише поточний стан, а й історію змін станів. У Tetris це могло б допомогти передбачати наслідки серій дій, однак RNN мають більшу обчислювальну складність і потребують обережного налаштування для стабільного тренування.

При побудові нейронної мережі для агента були враховані такі аспекти:

- Кількість шарів та нейронів: базова архітектура має три прихованих шари по 128 нейронів. Проте експериментально перевірялись варіанти з 64, 256 і навіть 512 нейронами на шар, що дозволяло підвищити виразність моделі, але вимагало більшого обсягу тренувальних даних та часу на навчання.
- Регуляризація: для запобігання перенавчанню (overfitting) доцільним є використання технік регуляризації, таких як dropout або L2-регуляризація. У даній реалізації проводились експерименти з dropout із ймовірністю 0.2–0.3, що дозволило покращити генералізацію політики агента.
- Batch Normalization: хоча у базовій архітектурі цей елемент відсутній, batch normalization міг би бути корисним для прискорення збіжності навчання та підвищення стійкості до зміни гіперпараметрів.

Однією з важливих властивостей побудованої нейронної мережі є можливість масштабування під різні цілі:

- Для роботи із більшими полями (наприклад, 20×20 замість стандартного 10×20) можна збільшити кількість шарів або нейронів.
- Для роботи з візуальними даними можна замінити приховані шари на згорткові.
- Для передбачення послідовностей дій у майбутньому можливо інтегрувати рекурентні компоненти (LSTM, GRU).

Під час тренування модель логуються з використанням TensorBoard, що дає змогу відстежувати:

- втрату функції (loss);

- середню винагороду;
- значення Q ;
- ефективність політики.

Це положення органічно вписується в загальну структуру виконаної роботи та сприяє глибшому аналізу якості навчання, полегшуючи підбір гіперпараметрів.

У результаті спостерігаються такі висновки з логів:

- Q -значення показують стабільне зростання — агент навчився оцінювати вигідні дії.
- Втрати (loss) зменшуються плавно — модель поступово покращує передбачення.
- Стратегія ϵ -зменшення працює: частка випадкових дій знижується від 100% до $\sim 1\%$.

3.3 Вибір інструментів та середовища розробки

Вибір програмного стеку — один із критично важливих етапів реалізації подібних систем. Враховуючи обмежені ресурси, було прийнято рішення відмовитись від складних фреймворків типу PyTorch або JAX і обрати Keras на базі TensorFlow через його доступність, велику кількість прикладів, швидкість розробки та зручність візуалізації.

Для реалізації системи було використано потужні інструменти Python, серед яких Keras, NumPy, OpenCV[16], TensorBoard, Jupyter Notebook та Google Colab. Такий набір інструментів забезпечує весь цикл розробки: від моделювання середовища до навчання агента та візуалізації результатів. Вибір цих бібліотек обумовлений їхньою широкою підтримкою, гнучкістю, продуктивністю та наявністю великої кількості прикладів і документації. Це дозволяє ефективно будувати та відлагоджувати систему навіть в умовах обмежених ресурсів.

Популярність Python зумовлена не лише лаконічним синтаксисом і низьким порогом входження, але й наявністю широкого спектра бібліотек, які охоплюють усі етапи розробки інтелектуальних систем: від збирання й обробки даних до побудови

моделей і їхнього розгортання в продуктивному середовищі. Python підтримується академічною і комерційною спільнотами, що гарантує швидкий розвиток інструментів і наявність великої кількості прикладів, документації та відкритих проєктів. У контексті розробки агента для гри в Tetris із використанням глибокого навчання з підкріпленням (Deep Reinforcement Learning) це дозволило сконцентруватися безпосередньо на побудові логіки навчання, уникаючи зайвих труднощів, пов'язаних із низькорівневою реалізацією.

З метою підвищення продуктивності системи також розглядався варіант прискорення критичних операцій через використання Numba — компілятора Python-функцій у машинний код. Це дозволило досягти пришвидшення деяких функцій обробки стану поля в 3–5 разів.

Для організації та керування експериментами використовувалась система збереження параметрів та результатів у форматі YAML/JSON. Це дало змогу відтворити експеримент при зміні лише кількох параметрів, що важливо для репрезентативності досліджень. Також реалізовано механізм фіксації псевдовипадкового зерна (random seed), що дозволяє повторюваність експериментів.

Варто підкреслити роль Google Colab не лише як засобу запуску, а і як платформи для спільного аналізу результатів — це забезпечило швидку зворотну реакцію під час навчального процесу. Застосування Google Colab особливо вигідне в академічному середовищі. Це дозволяє працювати з GPU без потреби витрат на фізичне обладнання, автоматично зберігати ноутбуки на Google Drive, а також ділитись експериментами з викладачами або колегами. Ще однією перевагою такого стеку є висока ступінь інтеграції між компонентами: Keras легко підключається до TensorBoard, NumPy природно вбудовується в усі рівні обробки даних, а OpenCV дозволяє одразу виводити результат кожного кроку гри. Ця взаємодія знижує бар'єри для експериментування та прискорює ітерації. Також без потреби у встановленні додаткового ПЗ.

Основні інструменти, використані в системі:

- Python 3.10 — остання стабільна версія мови на момент розробки, яка підтримує нові можливості (наприклад, структуровані шаблони match-case, покращене

типізування та нові модулі), що спрощують розробку великого коду та покращують читабельність. Завдяки цьому стало можливим структурувати логіку поведінки агента у зрозумілій та гнучкій формі, що особливо актуально в умовах багаторазових ітерацій моделі та складної логіки прийняття рішень.

- NumPy — базова бібліотека для наукових обчислень у Python, яка дозволяє ефективно працювати з векторами, матрицями та багатовимірними масивами. У межах проєкту вона застосовується для представлення ігрового поля у вигляді числової матриці, а також для обчислення станів, обчислення функцій винагороди та виконання трансформацій над простором станів. Наприклад, вектор стану може містити інформацію про висоту кожного стовпця, наявність порожніх клітин тощо — і всі ці обчислення зручно реалізуються через NumPy.
- Keras (з базою на TensorFlow) [13][14] — високорівнева бібліотека для глибокого навчання, яка дозволяє створювати та тренувати штучні нейронні мережі з мінімальними зусиллями. У цьому проєкті вона використовується для реалізації архітектури Q-мережі (Q-Network), яка наближає оптимальну функцію значення дії (Q-функцію) в контексті навчання з підкріпленням. Keras дозволяє швидко тестувати різні варіанти гіперпараметрів (кількість шарів, кількість нейронів, функції активації тощо), завдяки чому модель може бути адаптована до конкретних особливостей гри в Tetris.
- OpenCV — одна з найпотужніших бібліотек для комп'ютерного зору. У даному проєкті вона використовується для візуалізації ігрового процесу: відображення змін стану ігрового поля після кожної дії агента, запису відео та фреймів. Це дозволяє не лише аналізувати процес навчання, а й виявляти помилки у взаємодії агента з середовищем. OpenCV також може бути використана для інтеграції моделі з реальними зображеннями, наприклад, для розпізнавання гри на фізичному пристрої.
- TensorBoard — спеціалізований інструмент для моніторингу процесу тренування нейронної мережі. У межах розробки системи з підкріпленням він дозволяє відстежувати зміну функції втрат (loss function), середню винагороду

за епоху, кількість кроків до завершення гри та інші важливі метрики. Така візуалізація дає змогу оперативно коригувати архітектуру або гіперпараметри моделі на основі реальних результатів.

- Jupyter Notebook — інтерактивне середовище для розробки, яке дозволяє комбінувати код, текст, графіки та вивід результатів в єдиному документі. Воно стало ключовим інструментом під час етапу експериментального аналізу, оскільки дозволяє швидко перевіряти ефективність окремих частин коду, будувати графіки винагороди, змінювати структуру мережі, експериментувати з функцією винагороди або евристиками без потреби у повному перезапуску всієї програми.
- Google Colab — хмарне середовище, що забезпечує безкоштовний доступ до GPU (графічних процесорів), що є критично важливим для ефективного тренування глибоких нейронних мереж. У рамках цього проєкту Google Colab дозволив уникнути обмежень, пов'язаних із відсутністю високопродуктивного локального обладнання. Крім того, інтеграція з Google Drive забезпечила зручне збереження проміжних результатів та спільну роботу над проєктом.

Загальна оцінка інструментального вибору

Поєднання вищезазначених бібліотек і середовищ дозволяє реалізувати повний цикл розробки інтелектуального агента — від моделювання ігрового середовища до навчання та аналізу ефективності. Така інтеграція забезпечує не лише високу продуктивність, але й масштабованість, гнучкість та прозорість процесу. У контексті дослідження системи на базі гри Tetris, що характеризується складною динамікою, непередбачуваністю середовища та необхідністю прийняття рішень у реальному часі, саме використання Python та його екосистеми інструментів дозволило досягти якісного результату за обмежених ресурсів і термінів.

3.4 Тестування створеного інтелектуального агента для гри в Tetris у вигляді бота

У межах тестування важливо забезпечити не лише коректність роботи агента у стандартних сценаріях, а й оцінити його поведінку в нетипових ситуаціях. Для цього було використано методи «стрес-тестування», які включали:

- тренування на штучно зменшеному полі (наприклад, 6×10);
- зміна набору фігур (видалення деяких типів);
- підвищення швидкості падіння фігур.

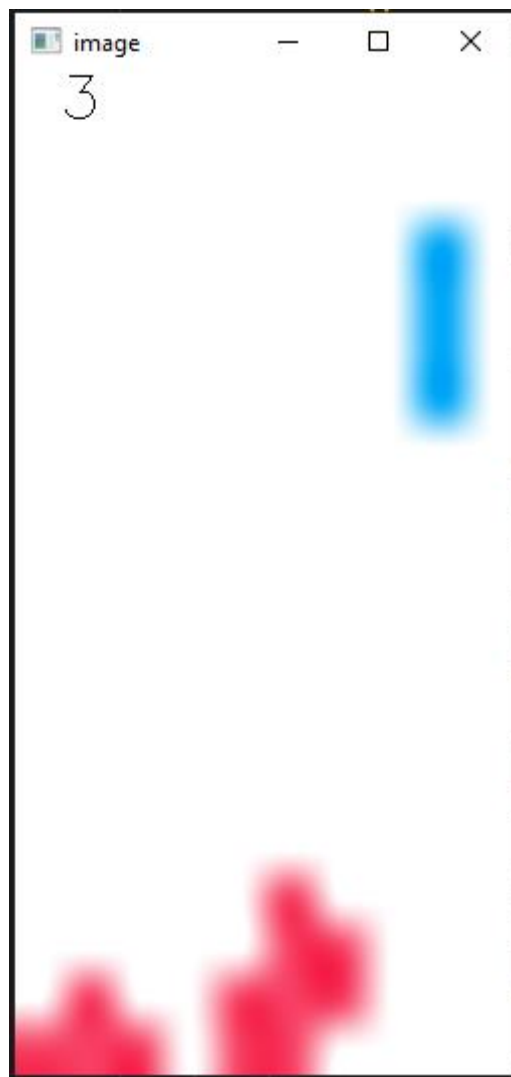


Рисунок 3.2 — Візуальний вигляд програми

Ці експерименти дозволили оцінити здатність агента до адаптації, стабільності політики та виявити обмеження поточної архітектури мережі. Наприклад, надто велике поле викликає труднощі у формуванні ефективної політики, якщо вхідний шар не масштабується відповідно до зростання кількості ознак.

Також було протестовано стратегію збереження контрольних точок (checkpoints), що дозволяє відновлювати навчання з місця зупинки. Це особливо важливо в умовах обмеженого часу або при використанні нестабільних серверів.

Навчання агента відбувалося з поступовою адаптацією параметрів, використовуючи ϵ -жадібну стратегію, буфер досвіду, оптимізатор Adam та функцію втрат MSE. Важливими аспектами було уникнення перенавчання, збереження різноманітності досвіду, контроль за динамікою функції втрат. Тестування проводилось на серії епізодів із фіксованими параметрами, що дозволило об'єктивно оцінити продуктивність агента. Було проведено порівняння з базовими стратегіями та візуалізовано результати через TensorBoard. Агент показав поступове покращення ефективності, що підтверджує валідність обраного підходу.

Важливою складовою успішного тренування агента є коректне налаштування гіперпараметрів. Нижче наведено їх перелік із поясненнями:

Таблиця 3.1 — Гіперпараметри

Параметр	Значення	Пояснення
Кількість епізодів	1000	Повні ігрові сесії для навчання
Розмір буфера пам'яті	20 000	Обсяг попереднього досвіду
Початкове ϵ (epsilon)	1.0	Імовірність випадкового вибору дії
Мінімальне ϵ	0.01	Нижня межа експлуатації
Швидкість зменшення ϵ	0.995	Регулює перехід від дослідження до експлуатації
Розмір мінібатчу	512	Кількість прикладів для одного кроку навчання
Кількість нейронів	[128, 128, 128]	Глибина та складність мережі
Функція втрат	MSE	Відповідає за обчислення помилки
Оптимізатор	Adam	Популярний адаптивний оптимізатор
Дисконтування (γ)	0.99	Врахування майбутніх винагород

Ці параметри були підібрані експериментально шляхом поступового збільшення складності моделі, що забезпечує системність і послідовність викладеного матеріалу.

Гіперпараметри відіграють вирішальну роль у навчанні агентів із використанням алгоритмів навчання з підкріпленням. У контексті гри Tetris, де динаміка середовища є складною, а простір дій та станів — надзвичайно великим, навіть незначна зміна значень гіперпараметрів може суттєво вплинути на кінцевий результат. Коефіцієнт ϵ у стратегії ϵ -greedy визначає баланс між дослідженням і експлуатацією: високе початкове значення (1.0) сприяє накопиченню досвіду, тоді як занадто повільне зниження ϵ затримує перехід до стабільних стратегій, а швидке зменшення може завадити виявленню оптимальних дій.

Швидкість навчання (α) також істотно впливає на конвергенцію: надто велика призводить до нестабільності та флуктуацій втрат, надто мала — до затримки в оновленнях. Для цієї задачі було визначено оптимальне значення $\alpha = 0.001$ в поєднанні з оптимізатором Adam, що забезпечує стабільне зменшення похибки при збереженні динаміки навчання. Розмір буфера досвіду впливає на різноманітність навчальних прикладів: буфер обсягом 50 000–100 000 переходів дозволяє накопичити достатню кількість унікальних ситуацій, водночас не перевантажуючи пам'ять застарілими даними. У ході експериментів було також відзначено, що розмір мінібатчу 64 забезпечує оптимальний компроміс між якістю оновлення ваг і швидкістю обчислень.

Після кожного тренувального циклу агент тестувався на декількох фіксованих середовищах з однаковими стартовими умовами. Основні метрики включали середню кількість балів за гру, кількість очищених ліній, число дій до завершення гри та середнє значення Q-функції. Тестування проводилося як із візуалізацією, так і в headless-режимі для пришвидшення оцінювання.

Розроблена система виявилася стабільною, розширюваною та придатною для подальших експериментів. Модульна архітектура дозволяє легко змінювати параметри середовища (наприклад, розмір поля), замінювати архітектуру нейромережі без впливу на інші компоненти, повторно використовувати буфер пам'яті в інших задачах та візуалізувати процес навчання для полегшення налагодження.

Серед основних проблем, із якими зіткнулися під час реалізації, — нестабільність навчання залежно від функції винагороди, необхідність ручного підбору гіперпараметрів, обмеження обчислювальних потужностей без доступу до GPU та схильність деяких моделей до перенавчання на певних типах фігур. Для подолання цих складнощів було протестовано кілька стратегій: зменшення глибини буфера пам'яті, заміну функції втрат, адаптивне налаштування ϵ , що впливало на стабільність і швидкість навчання. Наприклад, зменшення розміру мінібатчу призводило до більшої флуктуації результатів, але суттєво прискорювало обробку.

Тренування на GPU дозволило скоротити час одного епізоду в кілька разів, що дало змогу провести більше ітерацій та експериментів. Також для пришвидшення було використано headless-режим запуску гри. Усі експерименти проводилися в середовищі Google Colab, що забезпечило зручність масштабування та подальшого аналізу результатів.

3.5 Аналіз результатів роботи створеного інтелектуального агента для гри в Tetris

Результати експериментів свідчать про здатність агента ефективно навчатися на основі досвіду та формувати стратегії, що перевершують прості евристики. Виявлені обмеження, такі як залежність від функції винагороди, чутливість до гіперпараметрів, повільне навчання — є типовими для методів DRL і можуть бути подолані шляхом використання модифікованих підходів. Перспективним напрямом є застосування Double DQN, Dueling DQN, а також реалізація більш складних середовищ, включаючи 3D або мультиагентні варіанти гри. Застосування агентів RL у реальних задачах вимагає додаткової уваги до безпеки, валідації та прозорості алгоритмів, що підкреслює важливість системного підходу до розробки та тестування.

Аналіз результатів передбачає не лише оцінку абсолютних значень метрик, а й виявлення трендів у динаміці навчання. Наприклад, поступове зростання середньої винагороди протягом навчання свідчить про формування стійкої політики, тоді як флуктуації можуть вказувати на нестабільність моделі або неправильну функцію винагороди. Візуалізація за допомогою інструментів типу TensorBoard дозволяє не лише моніторити процес, а й вчасно виявляти потенційні проблеми, такі як переадаптація або зупинка покращення. Крім того, порівняння з іншими агентами (наприклад, випадковою або жадібною політикою) дозволяє дати кількісну оцінку

ефективності застосованого методу. Загалом, аналіз результатів є важливим інструментом у прийнятті рішень щодо подальшого вдосконалення системи.

Окрім стандартної реалізації з використанням OpenCV, для візуалізації гри можуть бути застосовані альтернативні програмні інструменти, що забезпечують ширший спектр функціональності, зручності інтеграції та адаптації до різних платформ. Одним із таких засобів є PyGame — популярна бібліотека для створення 2D-ігор, яка підтримує інтерактивні елементи та має вбудовані засоби для обробки графіки та подій. Завдяки простому API та великій кількості навчальних матеріалів, PyGame є зручним варіантом для швидкої побудови візуального представлення гри. У свою чергу, Unity у поєднанні з ML-Agents відкриває можливості для побудови високореалістичних 3D-середовищ із повноцінною фізикою та інтерфейсами, що дає змогу реалізувати глибшу симуляцію взаємодії агента з середовищем. Крім того, використання веб-інтерфейсів, таких як Flask або Streamlit, дозволяє демонструвати поведінку агента безпосередньо у браузері, що є корисним для інтеграції у навчальні або демонстраційні платформи та забезпечує швидкий доступ до системи з різних пристроїв.

Також було відзначено, що агенти схильні до формування короткозорих стратегій, які оптимізують миттєву винагороду, але не забезпечують довготривалої стабільності. У цьому контексті перспективними є підходи на кшталт багатокрокової винагороди (n-step TD) або використання цільової мережі з повільним оновленням.

Загалом, система продемонструвала здатність до навчання, ефективну інтеграцію компонентів і добрі результати в порівнянні з випадковою політикою. Проте для досягнення високого рівня генералізації та переносимості стратегій у майбутньому слід розглядати більш складні архітектури мереж, зокрема Dueling DQN або Actor-Critic, а також впроваджувати системи автооптимізації гіперпараметрів (наприклад, Optuna або Hyperopt).

З метою покращення взаємодії користувача з системою навчання агента доцільно реалізувати графічний інтерфейс користувача (GUI) для налаштування параметрів та керування процесом тренування. Такий інтерфейс може містити форму для введення гіперпараметрів (наприклад, ϵ — коефіцієнт жадібності, γ — коефіцієнт

дисконтного фактору, розмір буфера досвіду), кнопки запуску, зупинки та перезапуску тренування, а також блок візуалізації графіків, що оновлюються в реальному часі. Це дозволяє оперативно реагувати на зміни в навчанні, не вдаючись до ручного аналізу логів або скриптів. Реалізація такого інтерфейсу можлива як у вигляді десктоп-додатків за допомогою PyQt або Tkinter, так і у форматі інтерактивних веб-програм із використанням Streamlit, що особливо зручно для демонстрацій та колективної роботи.

Попри те, що дослідження виконується в академічному контексті, варто приділити увагу питанням безпеки при реалізації систем на основі навчання з підкріпленням. По-перше, важливо забезпечити ізоляцію середовища тренування від інших системних процесів, щоб уникнути конфліктів або втручання у виконання агентів. По-друге, необхідно реалізувати механізм збереження контрольних точок (checkpoints), що дозволяє не втратити прогрес у випадку помилок чи збоїв. Також варто валідувати всі вхідні дані, які подаються в систему, аби уникнути ситуацій, що призводять до неочікуваних або некоректних станів. Крім того, моніторинг споживання ресурсів є критично важливим — надмірне навантаження на CPU або GPU може не лише знизити продуктивність, а й пошкодити апаратне забезпечення. Застосування вищезазначених заходів є не лише технічно доцільним, а й критично важливим у разі планування реального застосування моделі.

Створена система має високий потенціал інтеграції з іншими платформами та системами. Зокрема, її можливо адаптувати для використання в навчальних ігрових платформах, де демонстрація роботи алгоритмів машинного навчання відбувається в інтерактивному форматі. Іншим перспективним напрямом є створення мобільних застосунків, які дозволяють користувачам експериментувати з агентами на власних пристроях. Крім того, система може бути адаптована до роботи у віртуальних або доповнених середовищах (VR/AR), що створює ще глибший досвід взаємодії з агентом. У поєднанні з хмарними обчислювальними платформами, такими як Google Cloud або AWS, система здатна виконувати тренування з великою кількістю параметрів або на значному обсязі даних.

Протягом реалізації дипломного проєкту було виконано повний цикл розробки та дослідження системи навчання агента для гри у Tetris. Процес включав аналіз ігрового середовища, підбір алгоритму навчання (DQN), створення середовища на Python, розробку агента з глибокою нейронною мережею, а також тренування, тестування та аналіз ефективності. У ході роботи було виявлено ряд типових проблем, таких як нестабільність процесу навчання, перенавчання, надто швидке зменшення параметра ϵ . Для їх усунення застосовувались сучасні практики: впроваджено буфер досвіду (replay buffer), адаптовано графік зміни ϵ , а також модифіковано архітектуру нейромережі з урахуванням специфіки гри.

Використання алгоритму Deep Q-Network (DQN) виявилось ефективним завдяки його здатності працювати з великими просторами станів, що є характерними для таких ігор, як Tetris. Завдяки модульній архітектурі було можливо швидко змінювати або оновлювати окремі компоненти системи без потреби повної перебудови коду. Важливою перевагою стала візуалізація навчання через TensorBoard, що дозволяла в реальному часі аналізувати динаміку змін показників. Агент продемонстрував високу адаптивність до зміни параметрів гри, таких як розміри поля або набір фігур, а також самостійно формував ефективну політику, що не базувалась на заздалегідь заданих шаблонах. Це свідчить про практичну ефективність застосованого підходу.

Водночас реалізована система має певні обмеження, зокрема повільне навчання, що потребує значної кількості епізодів, обмеженість функції винагороди та використання лише базової архітектури нейромережі. У поточній реалізації система працює на CPU, що ускладнює масштабування на великі обсяги даних або більш складні середовища. У подальшому можливе впровадження сучасних підходів, таких як Double DQN, Dueling DQN, Actor-Critic або мета-навчання, що покращить здатність агента до узагальнення та прискорить процес адаптації до нових середовищ.

У ході реалізації проєкту було застосовано ряд практик перевірки якості, зокрема модульне та інтеграційне тестування, порівняння з випадковою політикою (random policy), аналіз логів у TensorBoard, що дозволяло робити коригування параметрів у разі виявлення деградації навчання. Середньоарифметичний результат

DQN-агента був у 3–5 разів вищим, ніж у випадкового, що свідчить про успішність реалізованої стратегії. Було також оцінено якість агента за кількома метриками: середній рахунок, середня кількість очищених ліній, середнє значення Q-функції — зростання всіх цих показників корелювало з тривалістю тренування.

Аналіз проблем узагальнення політики показав, що в нових умовах агент іноді демонструє надмірну переадаптацію до середовищ, подібних до тренувальних. Це підкреслює важливість розширення простору даних для навчання, використання домен-рандомізації або ансамблевих стратегій. Перспективним також є застосування багатокрокової винагороди (*n-step returns*), що дозволяє агенту враховувати наслідки дій у довшій перспективі, та *transfer learning*, який зменшує потребу в повторному тренуванні.

Не можна оминути й етичні аспекти застосування навчання з підкріпленням. У реальних середовищах (медицина, автономний транспорт) необхідно враховувати ризики, пов'язані з неконтрольованими діями агента або дискримінаційними рішеннями. Хоча в рамках гри такі ризики мінімальні, обговорення цих питань є доцільним навіть у навчальному проєкті.

Розроблений агент може бути адаптований до інших ігор зі схожою логікою, таких як Snake (де складність гри зростає зі збільшенням довжини змії), Blockout (3D-варіант Tetris), або Puzzle Quest (гра з послідовним вибором дій). Така універсальність архітектури дозволяє повторно використовувати більшість компонентів системи.

У навчальному процесі система може слугувати як інструмент візуалізації роботи алгоритмів RL, основа для лабораторних робіт або змагань між студентами з побудови агентів. Також вона може бути корисною як базовий шаблон для розширення у рамках курсових чи дипломних проєктів.

3.6 Висновок до розділу 3

Реалізована система має чітку модульну структуру, яка дозволяє ефективно тренувати агента, оцінювати його продуктивність і демонструвати результати. Застосування сучасних бібліотек Python (Keras, NumPy, OpenCV) забезпечує

простоту та гнучкість розробки. Цей аспект також привертає увагу з огляду на його значущість у контексті обраної тематики.

Варто також зазначити, що гнучкість побудованої системи дозволяє використовувати її як шаблон для інших середовищ навчання з підкріпленням. Наприклад, при заміні логіки середовища з Tetris на іншу гру або задачу (керування об'єктом, ухилення від перешкод, оптимізація шляху), решта компонентів системи (агент, механізми тренування, логування) залишаються практично без змін. Таким чином, розроблена архітектура може бути використана як основа для подальших досліджень або створення навчальних систем із більшою складністю. Це особливо цінно в контексті сучасної тенденції до повторного використання моделей і скорочення вартості експериментів. Це положення органічно вписується в загальну структуру виконаної роботи.

ВИСНОВКИ

Дослідження підтвердило доцільність використання алгоритму DQN для навчання агента гри Tetris. Обрана гра забезпечує складне, динамічне середовище з високою варіативністю станів, що дозволяє ефективно тестувати алгоритми глибокого навчання з підкріпленням. Агент поступово набував стратегічної поведінки, що засвідчено зростанням основних метрик.

Використання буфера досвіду, ϵ -жадібної політики та нейронної мережі дозволило досягти стабільного покращення продуктивності агента. Разом з тим, виявлені обмеження — зокрема спрощена модель і відсутність вдосконалених модифікацій DQN — створюють підґрунтя для подальших досліджень та удосконалення підходу.

Розгляд поняття інтелектуального агента та його реалізація у вигляді DQN-агента для гри в Tetris дає змогу глибше зрозуміти принципи сучасного навчання з підкріпленням. Інтелектуальні агенти, наділені автономністю, адаптивністю та здатністю до навчання, виступають ефективним інструментом для вирішення задач у динамічних середовищах. В основі їх роботи лежать математичні моделі, зокрема марковські процеси прийняття рішень, що формалізують процес взаємодії агента з середовищем і забезпечують побудову оптимальної політики.

Застосування алгоритмів типу DQN дозволяє агенту опрацьовувати складні ситуації, оптимізувати дії та враховувати як короткострокові, так і довгострокові наслідки власних рішень. Водночас використання Tetris як середовища для навчання є показовим прикладом складності, яка потребує глибокого планування, ефективної оцінки станів і адаптації до випадкових подій. Представлені UML-діаграми структурують архітектуру рішення та демонструють тісну взаємодію між логікою гри та навчальним агентом, що в підсумку формує єдину адаптивну систему штучного інтелекту.

Реалізована система має чітку модульну структуру, яка дозволяє ефективно тренувати агента, оцінювати його продуктивність і демонструвати результати. Застосування сучасних бібліотек Python (Keras, NumPy, OpenCV) забезпечує

простоту та гнучкість розробки. Цей аспект також привертає увагу з огляду на його значущість у контексті обраної тематики.

Варто також зазначити, що гнучкість побудованої системи дозволяє використовувати її як шаблон для інших середовищ навчання з підкріпленням. Наприклад, при заміні логіки середовища з Tetris на іншу гру або задачу (керування об'єктом, ухилення від перешкод, оптимізація шляху), решта компонентів системи (агент, механізми тренування, логування) залишаються практично без змін. Таким чином, розроблена архітектура може бути використана як основа для подальших досліджень або створення навчальних систем із більшою складністю. Це особливо цінно в контексті сучасної тенденції до повторного використання моделей і скорочення вартості експериментів. Це положення органічно вписується в загальну структуру виконаної роботи.

СПИСОК ВИКОРИСТАНИХ ДЖЕРЕЛ

1. Бондаренко В. М., Бондаренко І. В. Штучний інтелект: теорія і практика. — Київ: Видавничий дім «Слово», 2018. — 392 с.
2. Морозов В. А. Основи машинного навчання та інтелектуальних систем. — Київ: Ліра-К, 2019. — 376 с.
3. Коваль О. В. Ігрове програмування. Створення комп'ютерних ігор засобами Python. — Харків: Фоліо, 2020. — 288 с.
4. Палій В. П., Гладкий Є. В. Основи теорії алгоритмів та штучного інтелекту. — Львів: Новий Світ, 2017. — 412 с.
5. Шевчук В. С. Марковські процеси прийняття рішень у завданнях штучного інтелекту. — Київ: Наукова думка, 2016. — 324 с.
6. Ткаченко А. Глибоке навчання та reinforcement learning: український досвід. — Штучний інтелект і машинне навчання, №2(14), 2021, с. 45–58.
7. Соловійов А. В. Моделювання інтелектуальних систем: навчальний посібник. — Харків: ХНУРЕ, 2018. — 250 с.
8. Кондратенко Ю. П. Нейронні мережі, нечіткі системи та еволюційні алгоритми. — Одеса: Чорномор'я, 2015. — 346 с.
9. Чала Н. Д. Сучасні методи машинного навчання в задачах обробки даних. — Київ: КНЕУ, 2020. — 298 с.
10. Власенко М. В. Комп'ютерні ігри: технології створення та моделювання. — Київ: КНУ імені Тараса Шевченка, 2019. — 280 с.
11. Tetris: історія гри. Українська правда. Життя [Електронний ресурс]. <https://life.pravda.com.ua/society/2021/06/6/244866/>
12. Python та штучний інтелект: приклади використання. DOU [Електронний ресурс]. <https://dou.ua/lenta/articles/python-ai-ukraine/>
13. Штучний інтелект та машинне навчання: огляд сучасних підходів. Наука і суспільство, №4, 2021. — с. 77–84.
14. Досвід використання нейронних мереж у комп'ютерних іграх. Вісник КНУ імені Тараса Шевченка. Серія «Кібернетика», №1, 2022. — с. 55–63.

15. Tetris як модель для когнітивних досліджень. Психологічний журнал, №3, 2020. — с. 101–110.
16. OpenCV та Python у створенні систем комп'ютерного зору. Комп'ютерні науки та інформаційні технології, №2, 2021. — с. 33–42.
17. Глибоке навчання з підкріпленням для гри Tetris: приклад з української практики. Штучний інтелект і робототехніка, №3, 2022. — с. 89–96.
18. Марковські процеси та їх застосування в управлінні. — Харків: ХНУ, 2017. — 310 с.
19. Reinforcement Learning та його застосування в реальних завданнях. Системи управління, навігації та зв'язку, №4, 2021. — с. 58–65.
20. Досвід використання TensorFlow та Keras в українських проєктах. DOU [Електронний ресурс]. <https://dou.ua/lenta/articles/tensorflow-keras-ukraine/>

ДОДАТОК А (обов'язковий) Протокол перевірки кваліфікаційної роботи

Назва роботи: Розробка бота для гри в «Тетріс»

Тип роботи: бакалаврська кваліфікаційна робота
(бакалаврська кваліфікаційна робота / магістерська кваліфікаційна робота)

Підрозділ кафедра комп'ютерних наук
(кафедра, факультет, навчальна група)

Коефіцієнт подібності текстових запозичень, виявлених у роботі
системою StrikePlagiarism 0.62 %

Висновок щодо перевірки кваліфікаційної роботи (відмітити потрібне)

- ✓ Запозичення, виявлені у роботі, є законними і не містять ознак плагіату, фабрикації, фальсифікації. Роботу прийняти до захисту
- ☐ У роботі не виявлено ознак плагіату, фабрикації, фальсифікації, але надмірна кількість текстових запозичень та/або наявність типових розрахунків не дозволяють прийняти рішення про оригінальність та самостійність її виконання. Роботу направити на доопрацювання.
- ☐ У роботі виявлено ознаки плагіату та/або текстових маніпуляцій як спроб укриття плагіату, фабрикації, фальсифікації, що суперечить вимогам законодавства та нормам академічної доброчесності. Робота до захисту не приймається.

Експертна комісія:

Яровий А.А., зав. каф. КН

(прізвище, ініціали, посада)

Колесницький О.К., проф. каф КН

(прізвище, ініціали, посада)

(підпис)

(підпис)

Особа, відповідальна за перевірку

Озеранський В.С.

(підпис)

(прізвище, ініціали)

З висновком експертної комісії ознайомлений(-на)

Керівник

(підпис)

Гармаш В.В.

(прізвище, ініціали, посада)

Здобувач

(підпис)

Волинець І.Ю.

(прізвище, ініціали)

ДОДАТОК Б (обов'язковий)

Лістинг програми

```

from typing import ValuesView, List, Optional
from keras import Model
from keras.models import Sequential
from keras.layers import Dense
from collections import deque
import numpy as np
import random

class DQNAgent:
    def __init__(self, state_size, mem_size=10000, discount=0.95,
                  epsilon=1, epsilon_min=0, epsilon_stop_episode=500,
                  n_neurons=(32, 32), activations=('relu', 'relu', 'linear'),
                  loss='mse', optimizer='adam', replay_start_size=None):
        assert len(activations) == len(n_neurons) + 1
        self.state_size = state_size
        self.memory = deque(maxlen=mem_size)
        self.discount = discount
        self.epsilon = epsilon
        self.epsilon_min = epsilon_min
        self.epsilon_decay = (self.epsilon - self.epsilon_min) / epsilon_stop_episode
        self.n_neurons = n_neurons
        self.activations = activations
        self.loss = loss
        self.optimizer = optimizer
        if not replay_start_size:
            replay_start_size = mem_size / 2
        self.replay_start_size = replay_start_size
        self.model = self._build_model()
    def _build_model(self) -> Model:

```

```

model = Sequential()
model.add(Dense(self.n_neurons[0], input_dim=self.state_size,
activation=self.activations[0]))
for i in range(1, len(self.n_neurons)):
    model.add(Dense(self.n_neurons[i], activation=self.activations[i]))
model.add(Dense(1, activation=self.activations[-1]))
model.compile(loss=self.loss, optimizer=self.optimizer)
return model

def add_to_memory(self, current_state, next_state, reward, done):
    self.memory.append((current_state, next_state, reward, done))

def random_value(self):
    return random.random()

def predict_value(self, state: np.ndarray) -> float:
    return self.model.predict(state)[0]

def act(self, state):
    state = np.reshape(state, [1, self.state_size])
    if random.random() <= self.epsilon:
        return self.random_value()
    else:
        return self.predict_value(state)

def best_state(self, states: ValuesView[List[int]]) -> List[int]:
    if random.random() <= self.epsilon:
        return random.choice(list(states))
    else:
        max_value: Optional[float] = None
        best_state: Optional[List[int]] = None
        for state in states:
            value = self.predict_value(np.reshape(state, [1, self.state_size]))
            if not max_value or value > max_value:
                max_value = value

```

```

        best_state = state

    return best_state

def train(self, batch_size=32, epochs=3):
    n = len(self.memory)
    if n >= self.replay_start_size and n >= batch_size:
        batch = random.sample(self.memory, batch_size)
        next_states = np.array([x[1] for x in batch])
        next_qs = [x[0] for x in self.model.predict(next_states)]
        x = []
        y = []
        for i, (state, _, reward, done) in enumerate(batch):
            if not done:
                new_q = reward + self.discount * next_qs[i]
            else:
                new_q = reward
            x.append(state)
            y.append(new_q)
        self.model.fit(np.array(x), np.array(y), batch_size=batch_size, epochs=epochs,
verbose=0)
        if self.epsilon > self.epsilon_min:
            self.epsilon -= self.epsilon_decay

```

ДОДАТОК В (обов'язковий)**ІЛЮСТРАТИВНА ЧАСТИНА****«РОЗРОБКА БОТА ДЛЯ ГРИ «TETRIS»»**

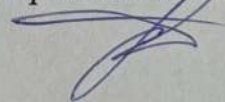
Виконав: студент 4 курсу, групи 4КН-216
спеціальності 122 – Комп'ютерні науки
(шифр і назва напрямку підготовки, спеціальності)



Волинець І. Ю.

(прізвище та ініціали)

Керівник: к.т.н., доцент кафедри АІТ



Гармаш В. В.

(прізвище та ініціали)

«___» _____ 2025 р.

03.06.2025р.



Рисунок В.1 — Візуальний вигляд програми

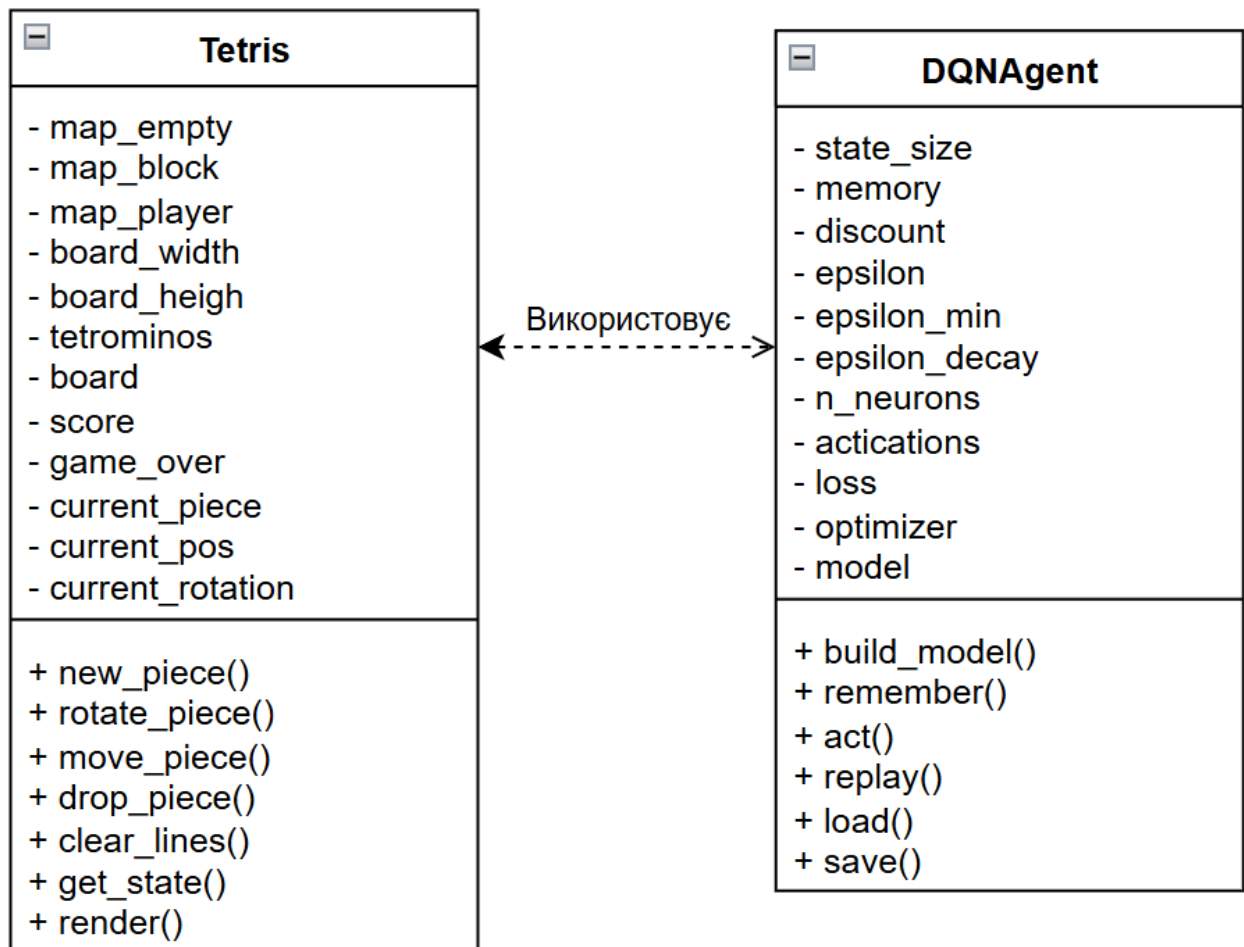


Рисунок В.2 — UML діаграма класів



Рисунок В.3 — Схема алгоритму роботи інтелектуального агента

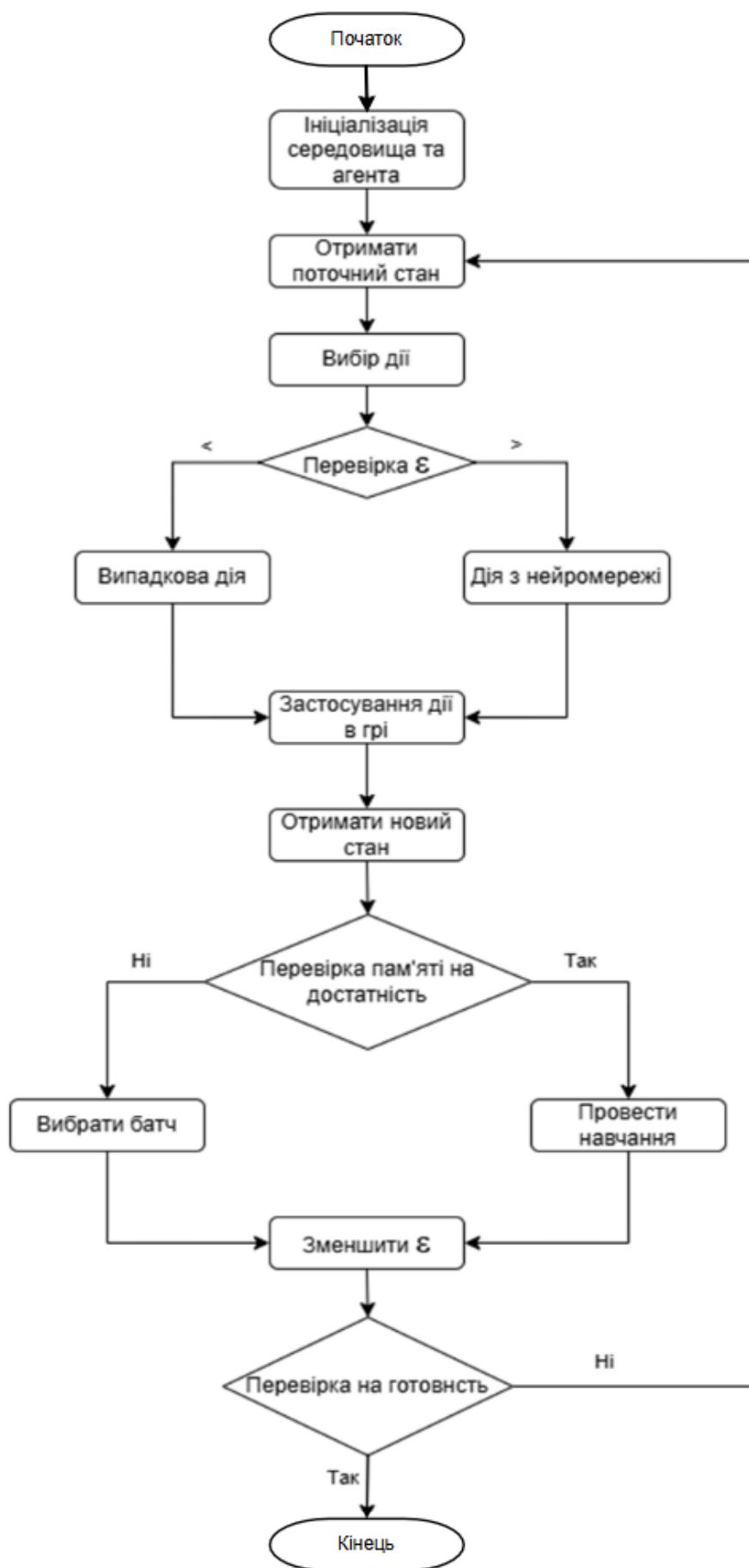


Рисунок В.4 — Схема алгоритму роботи бота

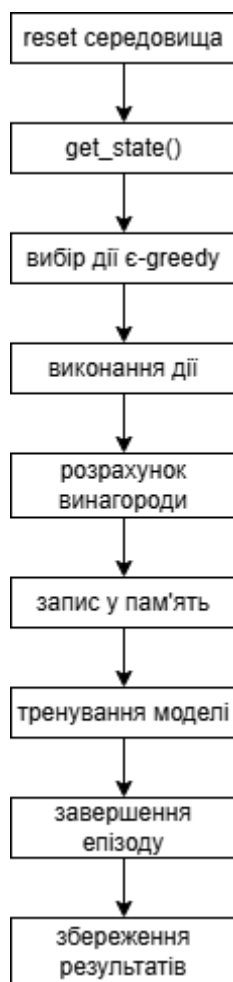


Рисунок В.5 — Схема життєвого циклу епізоду

ДОДАТОК Г Інструкція користувача

Опис програми

Ця програма реалізує інтелектуального агента, що навчається грати у гру Tetris за допомогою алгоритму Deep Q-Learning. Вона дозволяє:

- самостійно грати в Тетріс (інтерактивний режим),
- тренувати агента на великій кількості ігор,
- оцінювати продуктивність агента після тренування.

Вимоги до системи

Платформа:

- Linux або Windows з Python 3.7

Необхідні компоненти:

- Python 3.7
- TensorFlow (підтримка GPU рекомендована)
- Keras, Pillow, tqdm, OpenCV, tensorboard

Встановлення залежностей:

Варіант 1: Вручну

```
sudo apt install libgtk2.0-dev pkg-config
conda create --name py37tf python=3.7
conda install -c conda-forge opencv=4.1.0
conda install tensorflow-gpu keras pillow tqdm tensorboard
conda activate py37tf
```

Варіант 2: З файлу

```
conda create --name py37tf --file requirements.txt
conda activate py37tf
```

Режими використання

1. Інтерактивний режим (грати вручну)

Запускає класичну гру Tetris, де користувач керує фігурами з клавіатури.

```
python run_interactive.py
```

2. Режим тренування (навчання агента)

Починається процес навчання агента з нуля.

```
python run_train.py
```

3. Режим оцінювання (перевірка навченого агента)

Запускає тестування навченого агента:

```
python run_eval.py
```

Структура проєкту

- tetris.py – логіка гри Tetris
- dqn_agent.py – реалізація агента з DQN
- run_train.py – скрипт для тренування агента
- run_eval.py – скрипт для оцінки агента